

■ POUR LA SCIENCE

Février 2008

Édition française de Scientific American

Repenser l'écologie

**L'homme,
grand oublié
des stratégies
de conservation**



- Les sursauts gamma : le chant du cygne des étoiles
- Lutter contre la dégénérescence des neurones moteurs
- Des matériaux frustrés : les glaces de spin

M 02687 - 364 - F: 5,95 €



VISIT...

LANZAROTE
Caliente.COM

écoutez-vous

recherche
équipe
international
échanges

Multimédia, réseaux, cryptographie, robotique, réalité virtuelle, web...
autant de domaines dans lesquels plus que jamais la recherche
avance.

En **2008**, l'INRIA propose **150 profils**
de post-doctorats dans des
domaines scientifiques de pointe.

Jeunes scientifiques

Vous êtes titulaire d'une thèse dans le domaine de l'informatique
ou de l'automatique. Et vous souhaitez la compléter d'une expérience
au sein d'équipes scientifiques reconnues au niveau international pour
leur savoir-faire et la qualité de leurs travaux.

Venez faire votre post-doctorat à l'INRIA, un des principaux instituts de
recherche en France.

Rejoignez l'INRIA (INRIA et partenaires) et bénéficiez de l'enthousiasme qui
anime les 2 900 scientifiques et des moyens de pointe mis à leur disposition,
pour accroître votre expertise.

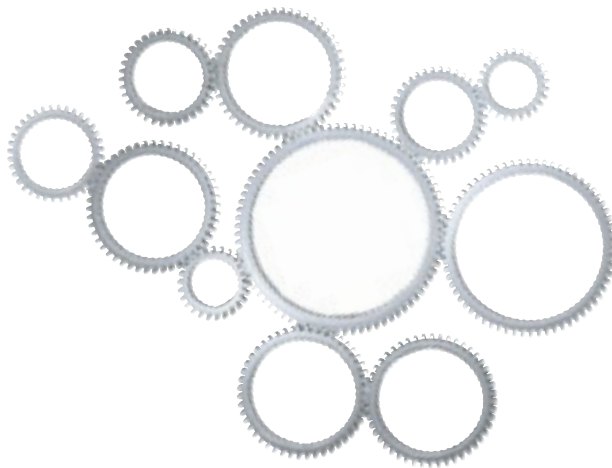
Ça commence ici et maintenant www.inria.fr/travailler

réseaux
sécurité
des
logiciels
systèmes
complexes
simulation
réalité
virtuelle
modélisation
du vivant



INSTITUT NATIONAL
DE RECHERCHE
EN INFORMATIQUE
ET EN AUTOMATIQUE





Anticiper : un art difficile

Savoir pour prévoir, afin de pouvoir.
Auguste Comte

Anticiper est une qualité essentielle de l'esprit. C'est entrevoir sans raisonnement explicite des idées ou des événements qui se réaliseront ultérieurement – si l'anticipation était pertinente ! L'anticipation est préconception, prénotion. Selon Kant, c'est une forme de connaissance *a priori*.

Dans quelque domaine que ce soit, on cherche à anticiper les conséquences de ses actes pour les ajuster aux contraintes. En matière de préservation de l'environnement, les résultats des mesures prises montrent combien l'anticipation est un art difficile. Ainsi, un anti-inflammatoire administré à des vaches a eu pour conséquence la quasi-disparition de trois espèces de vautours et la recrudescence de la rage en Asie ; l'extinction des loups dans l'Est des États-Unis a entraîné une augmentation du nombre des personnes atteintes de la maladie de Lyme... car les daims (et leurs tiques responsables de la maladie) se sont multipliés sans frein, leurs principaux prédateurs ayant été exterminés. En matière de gestion des populations animales, l'anticipation est souvent mise en défaut, en raison de la multitude des acteurs (voir *Repenser l'écologie*, page 38).

Manque d'anticipation aussi dans notre gestion des banques de sang de cordon ombilical. On a découvert, il y a une vingtaine d'années, que ce sang contient de précieuses cellules souches, capables de reconstituer tout le stock des cellules du sang. Prenant rapidement conscience que les cordons ombilicaux, jetés jusqu'alors comme des

rebutis opératoires, ont une valeur inestimable, certains pays ont organisé des banques pour les conserver. Le sang de cordon ombilical est aujourd'hui utilisé avec succès dans diverses greffes. Mais la France, qui pratique de nombreuses greffes, doit – faute d'avoir anticipé le potentiel de la méthode – importer 70 pour cent des sangs de cordons qu'elle utilise (voir *Comment stocker le sang de cordon ombilical ?*, page 26).

Ce sang est généralement destiné à soigner un proche ou un inconnu compatible, mais il peut aussi être stocké par les parents pour l'enfant qui vient de naître, telle une assurance biologique en cas de maladie ultérieure... Anticipation poussée à son paroxysme !

L'anticipation est une forme d'intelligence intuitive : on adapte son comportement à celui de l'autre. Ainsi, l'utilisation de scellés traduit la méfiance de l'homme pour ses congénères depuis... des millénaires. En témoignent des scellés retrouvés sur une tablette cunéiforme de quelque 3 700 ans. Les scellés ne servent pas à résister aux effractions, mais seulement à indiquer qu'effraction il y a eu (voir *Témoins d'effraction : les scellés*, page 82).

En matière d'intelligence, diverses études réalisées dans plusieurs pays ont montré que le quotient intellectuel aurait augmenté durant les 50 dernières années. Diverses explications sont proposées par les psychologues, qui toutes confirment qu'il est difficile de mesurer l'intelligence. L'intelligence globale continuera-t-elle à augmenter ou découvrira-t-on que les tests utilisés aujourd'hui sont inadaptés ? On peut anticiper sans trop risquer de se tromper que la seconde hypothèse est la bonne !

Françoise PÉTRY

février 2008

POUR LA SCIENCE N° 364

Science & actualité

5 Bloc-Notes

Didier Nordon

6 Science & gastronomie

Mangeons moins salé...

Hervé This

8 Science & économie

Quand les riches maigrissent...

Ivar Ekeland

10 Tribune des lecteurs

22 Les limites de la science

Les limites

de la physiologie humaine

Jean-Paul Richalet

26 Questions ouvertes

Comment stocker le sang de cordon ombilical ?

Éliane Gluckman

Le potentiel des cellules souches contenues dans le sang de cordon ombilical est immense. Ce sang est stocké soit dans des banques publiques de qualité contrôlée et faisant partie d'un réseau international, soit dans des banques privées à but lucratif. Comment concilier ces deux types de stockage ?

30 Point de vue

Ovnis : des enquêtes sujettes à caution

David Rossoni, Éric Maillot
et Éric Déguillaume

32 Présence de l'histoire

Kelvin, Perry et l'âge de la Terre

Philip England, Peter Molnar
et Frank Richter

12 Perspectives scientifiques

La tectonique intermittente • Réparer les muscles des myopathes ? • Climats tempérés sur planètes excentriques • Parité chromosomique • L'ancêtre terrestre des baleines • La Lune rajeunie • Effet Casimir dans un liquide • Des canaux membranaires à ouverture « électronique » • Le trésor du « bout du monde » • Évacuer une foule au plus vite • Des singes doués en calcul • Des spermatozoïdes invisibles • Crotale, parfum pour écureuils • Le Jupiter de Doliché... • Le mystère des nuages nacrés • Terminus à puces • Mars : l'effet de serre sent le soufre • Des lucioles moins efficaces • Cerveau de plongeur •

ÉCOLOGIE

→ 38

Repenser l'écologie

Peter Kareiva et Michelle Marvier

Il est illusoire, voire contre-productif, de préserver un écosystème isolé du reste du monde. Replaçant l'homme au centre de leur stratégie, les écologues repensent les interactions de tous les acteurs de la biodiversité.

ASTRONOMIE

→ 46

Les sursauts gamma, témoins de la mort violente des étoiles

Robert Mochkovitch

L'observation détaillée des sursauts gamma juste après leur détection a permis de découvrir les objets qui en sont à l'origine, et de préciser les mécanismes qui produisent ces bouffées de rayons gamma visibles depuis le fond de l'Univers.



NEUROBIOLOGIE

→ 54

Éviter la mort des neurones moteurs

Patrick Aebischer et Ann Kato

Dans la maladie de Charcot ou sclérose latérale amyotrophique, les neurones moteurs qui innervent les muscles dégénèrent progressivement : les patients finissent par être paralysés. Les neurobiologistes précisent les mécanismes de cette maladie et plusieurs traitements sont à l'étude.



Liquides et glaces de spin

Rafik Ballou et Claudine Lacroix

Dans certains réseaux cristallins, les énergies d'interaction entre aimants atomiques ou moléculaires ne peuvent être minimales toutes à la fois. Cette « frustration » conduit à des propriétés inhabituelles, que les physiciens cherchent à vérifier sur des matériaux réels. Avec quelques succès.

La redoutable efficacité des armes romaines

Yann Le Bohec

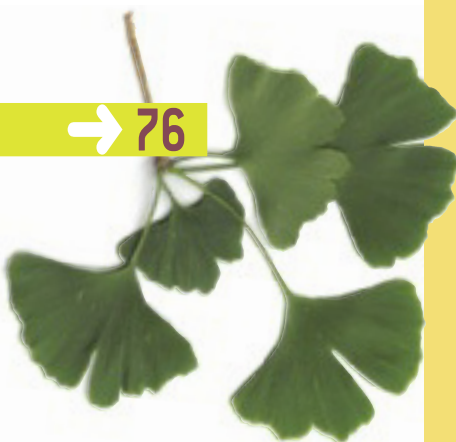
Depuis le XIX^e siècle, archéologues et amateurs multiplient les expériences pour reproduire et tester les armements romains. Les résultats attestent d'un grand savoir-faire, produisant les armes défensives ou offensives qui ont fait la puissance de Rome.



Ginkgo biloba : le rescapé et son algue

Jocelyne Trémouillaux-Guiller

L'arbre a traversé les millénaires presque inchangé. C'est un organisme d'exception, tant par son mode de reproduction que par l'algue que ses cellules hébergent.



Témoins d'effraction : les scellés

Roger Johnston

Utilisés depuis des millénaires, les scellés remplissent de multiples fonctions de sécurité. Mais de nouveaux développements devraient améliorer leur efficacité, aujourd'hui assez médiocre.



Visions scientifiques

90

Logique & calcul

Pierre, feuille, ciseaux...

Jean-Paul Delahaye

Joué comme un jeu de conquête, l'ancien jeu chinois de la pierre, de la feuille et des ciseaux éclaire les phénomènes de synchronisation et de maintien de la diversité animale.

96

Art & science

La petite sirène d'Ambroise Paré

Philippe Charlier

98

Idées de physique

Pont de Tacoma : la contre-enquête

Jean-Michel Courty et Édouard Kierlik

100 Analyses de livres

Bernard Heuvelmans, Jean-Jacques Barloy, *Les félins encore inconnus d'Afrique*, Bernard Heuvelmans •

L'Atlas des origines de l'homme, Douglas Palmer •

L'obsidienne, Laurent-Jacques Costa • *Volcans*,

Olivier Grunewald et Jacques-Marie Bardintzeff •

À la conquête de l'espace, Jacques Villain •

Les démons de Gödel, Pierre Cassou-Noguès •

Retrouvez des compléments d'information sur le site de *Pour la Science*

www.pourlascience.com

Sur la totalité des numéros : deux encarts d'abonnement pages 24 et 25, encarts commande de livres et abonnement pages 80 et 81.
En couverture : © Engrenages de Samuel et Pedro Velasco, *5W Infographics*.
Coureur : Rubberball Productions/Getty Images. Vignettes intégrées : Shutterstock/E. Isselée, I. Tischenko, Fouquin, J. Lye, C. Clapper, D. Zidar, D. Lundin, C. Erpenbeck, C. Musat, E. Buchko.

2008

Année Internationale
de la **Planète Terre**

www.anneeplaneteterre.com

La valeur n'attend pas le nombre des... citations

Comment déterminer la valeur d'un chercheur scientifique ? Telle est la question – du moins, aux yeux de tout ce que notre société compte d'« évaluateurs ».

En nos temps d'informatique reine, il est devenu facile de compter le nombre de fois où un chercheur est cité. Les évaluateurs recourent donc à cette opération, nommée « bibliométrie ». Mais ils se fondent sur une hypothèse douteuse, à savoir que la valeur d'un chercheur se mesure au nombre de fois où il est cité. Or on rapporte des exemples d'articles qui ont été souvent mentionnés parce qu'ils contenaient une erreur et que des auteurs ultérieurs tenaient à la pointer ! En tout cas, les résultats qu'obtient la bibliométrie attirent l'attention par leurs bizarreries. Des chercheurs au prestige médiocre se retrouvent classés devant des lauréats Nobel ou Fields. Voilà qui place les évaluateurs face à un dilemme dont on est curieux de voir comment ils vont sortir. Déclarer que les Nobel mal classés ont volé leur prix heurterait le sens des convenances et le respect des hiérarchies établies qui caractérisent les évaluateurs. Mais reconnaître que le classement par bibliométrie est peu pertinent exige-

rait qu'ils renoncent à la méthode vers laquelle va leur dévotion.

Le jugement porté sur un scientifique par ses pairs n'est pas fiable, tant la profession est infestée de coteries. Jouet des manipulations médiatiques, la célébrité d'un chercheur auprès du grand public n'est pas non plus preuve de sa valeur scientifique. La seule quasi-certitude est que, quelle que soit la hiérarchie établie entre les chercheurs par une époque, elle sera chahutée, sinon oubliée, par l'époque suivante.

Il faut s'y résigner. Apprécier la valeur d'un chercheur scientifique ne relève pas de la science, ni même de l'objectivité. De ce point de vue, les sciences sont dans la même situation que les arts. Voilà des milliers d'années qu'on est incapable de hiérarchiser les artistes. Cela n'a pas tué les arts, cela ne tuera pas les sciences.

Après la bibliométrie : l'Internetométrie !



Cercle Narcissique de Réflexion sur Soi

Quand on est jeune, on est « une force qui va ». Quand on s'interroge trop sur son passé, c'est qu'un ressort s'est détendu, sinon brisé. On a mal vieilli.

Que, depuis quelques années, le CNRS publie une revue portant sur l'histoire du CNRS est donc peut-être mauvais signe. Outre le narcissisme d'une institution dont certains membres se consacrent à l'histoire de ladite institution, il est bien possible que cette publication trahisse un désarroi. Les chercheurs sentent que leur activité a perdu de son prestige, et ils se consolent en se penchant sur les grandes luttes de leurs aînés. Malheureusement, évoquer un passé glorieux est une chose, contribuer à faire émerger de nouvelles heures glorieuses en est une tout autre. Peut-être, alors, les historiens du futur retiendront-ils comme indice de déclin d'une institution de recherche le fait qu'elle prenne sa propre histoire comme sujet de recherche.

Médiations à contre-emploi

Dans la vie familiale ou sociale, certains incidents mineurs se tassent sans trop de casse grâce au fait que les personnes qui auraient pu se sentir agressées ont fermé les yeux, ou pris les choses du bon côté, bref, ont toléré sans exiger de contrepartie une petite entorse à ce qu'elles auraient été en droit d'attendre. Elles ont ainsi évité que les choses s'enveniment. Ce n'est pas forcément

une démission. On incite parfois plus un malappris à rentrer dans le rang en négligeant son insolence qu'en en faisant trop de cas.

Il n'est donc pas certain que multiplier les médiateurs chargés de régler les différends aide à une meilleure entente entre les gens. Pour recourir à un médiateur, il faut qu'au moins un des protagonistes prenne mal l'événement qui s'est produit. Des cas qui auraient pu être vécus comme des contrariétés sans gravité se cristallisent et sont formalisés comme des incidents exigeant qu'on les règle. Les statistiques des « incivilités » répertoriées augmentent alors. Du coup, la nervosité ambiante aussi.

Notre époque pense que formaliser les problèmes est un pas vers leur solution. En fait, il n'en va pas toujours ainsi. De « formaliser » à « se formaliser », il n'y a pas loin. Parfois, mieux vaut s'abstenir de l'un comme de l'autre.

Médiateur en recherche d'emploi



Quand l'avenir a tort

Si, au temps du général de Gaulle, j'avais fait la surprenante prévision selon laquelle, un jour, le président de la République se montrerait en culottes courtes devant les photographes, on m'aurait traité d'aimable farceur. Cependant, l'avenir m'aurait donné raison. Par contre, si j'avais dit que ce serait un progrès des mœurs, l'avenir ne m'aurait donné ni raison ni tort. De même, si, vers 1960, j'avais dit qu'un jour, beaucoup de gens pourraient, de chez eux, accéder à une bibliothèque universelle, l'avenir m'aurait donné raison. Alors que si j'avais dit qu'il fallait œuvrer à ce que les gens puissent accéder à une bibliothèque universelle, l'avenir ne m'aurait donné ni raison ni tort.

La suite des événements donne raison à ceux qui l'ont prévue telle qu'elle a été, mais pas à ceux qui l'ont voulue telle qu'elle a été. En aucun cas, elle ne prouve que ce qui a été vaut mieux que ce qui aurait pu être. Si la connexion Internet devient un jour aussi indispensable que l'eau et le gaz à tous les étages (comme l'indiquaient encore certains immeubles de mon enfance), ce sera un fait, pas un jugement de valeur.

Donner tort à ceux qui ont perdu, c'est écrire l'histoire du point de vue des vainqueurs. Or écrire l'histoire du point de vue des vainqueurs donne l'agréable impression qu'on fait partie des vainqueurs. C'est pourquoi la propension à assimiler ceux qui ont gagné et ceux qui avaient raison se perpétuera longtemps. Et si l'avenir dément la prévision que je fais là, eh bien, c'est lui qui aura tort !

Mangeons moins salé...

... et moins gras en comprenant mieux la physico-chimie de la confection et du salage des savoureuses frites.



La friture est une opération merveilleuse qu'une certaine cuisine hygiéniste menace : nous craignons l'acrylamide, cancérigène, qui se forme lors des cuissons de produits contenant de l'amidon, telles les pommes de terre ; nous craignons la matière grasse, surtout quand elle a chauffé et que des réactions chimiques ont alors engendré des molécules aussi toxiques que l'acroléine formée quand l'huile fume ; et nous craignons le sel parce que, sans possibilité de régler la quantité de sel à l'intérieur du produit, la cuisinière ou le cuisinier saupoudre abondamment la surface des frites au risque de favoriser les maladies cardio-vasculaires des amateurs.

Pourtant que de beaux produits frits, des rissoles aux petits poissons, en passant par les beignets, les merveilles et, bien sûr, les frites ! La friture, en effet, a l'avantage de porter la surface de l'objet frit à une température bien supérieure à la température d'ébullition de l'eau, ce qui engendre une croûte sèche, dure, qui fait contraste avec l'intérieur, encore humide, tendre, moelleux. Or nos systèmes sensoriels reconnaissent surtout les contrastes : c'est pour cette raison que nous ne sentons plus le tabac dans une pièce enfumée après quelques minutes, ou que, dans les tunnels éclairés par de la lumière jaune, nous ne percevons la couleur des voitures qu'après quelques instants d'adaptation (au début tout paraît jaune). En cuisine, un aliment homogène est peu intéressant, et les cuisinières et cuisiniers gagnent à produire des contrastes.

La friture, d'autre part, produit une jolie couleur blonde, à cause des réactions chimiques entre les composés de surface qui, au lieu d'être dilués dans l'eau que renferment les aliments, sont concentrés et chauffés. Lors des fritures, prennent place les réactions de Maillard, entre des sucres (apportés par l'amidon, qui est composé de polymères du glucose) et par des acides aminés (libres ou dans les protéines), mais aussi bien d'autres réactions, d'oxydation, d'hydrolyse, de pyrolyse... Reste que les gourmands se contenteraient du meilleur : ils voudraient des frites croustillantes, dorées... mais sans graisse ni trop de sel.

Des mesures de la pression dans les frites donnent une clé de cette ambition gustative. La pression augmente progressivement, au cours de la friture, parce que les grains d'amidon de la surface commencent à s'empeser, soit dans la pâte à friture qui couvre l'aliment, soit dans l'intérieur des

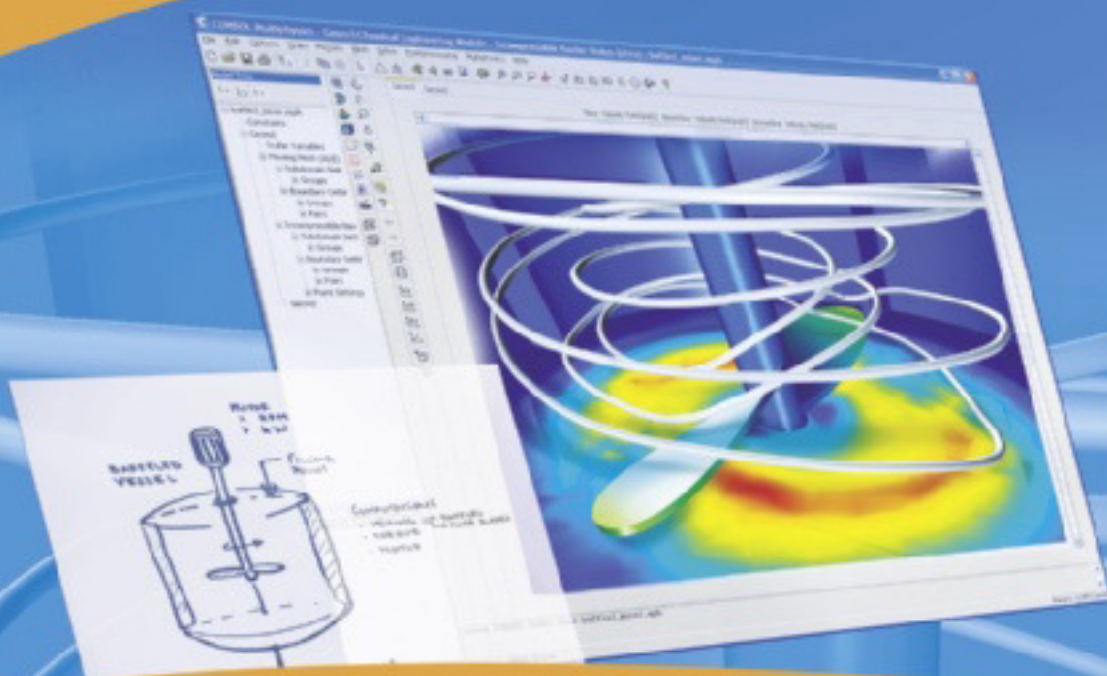
cellules des féculents tels que la pomme de terre. Gonflant par absorption de l'eau environnante, les grains d'amidon se soudent, puis, lorsque l'eau est évaporée, ils forment une croûte qui emprisonne la vapeur d'eau formée dans l'intérieur de l'aliment. Ce phénomène est particulièrement marquant lors de la confection des pommes de terre soufflées : lorsque les rondelles de pomme de terre sont plongées dans un bain d'huile très chaud, une couche de surface se met à boursoufler tant l'afflux de vapeur est rapide.

Pour des fritures plus classiques, le départ de l'eau de l'aliment, sous la forme de vapeur, par les orifices de la croûte, conduit à former des cavités pleines de vapeur. Lors du refroidissement des objets frits, cette vapeur se condense, ce qui s'accompagne d'une diminution de la pression... et l'huile de la surface est finalement réabsorbée ! Or plus la surface est rugueuse, plus elle « accroche » d'huile par des forces de surface. Pas étonnant, alors, que l'on ait mesuré que la rugosité des frites s'accompagne d'une teneur en huile supérieure, quand aucune mesure n'est prise pour éviter l'absorption au cours du refroidissement. N'oublions donc pas d'éponger les frites au sortir du bain, avant que la vapeur d'eau se recondense. Il en va de plusieurs dizaines de grammes d'huile par 100 grammes de pommes de terre frites !

Et le sel ? Il a bien peu de raison d'adhérer à la surface des frites, car il est hydrophile, alors que la surface des frites est couverte d'huile, donc hydrophobe. De fait, il a été mesuré récemment (V. Buck et S. Barringer, *Factors dominating adhesion of NaCl onto potato chips*, *Journal of Food Science*, vol. 72 (8), pp. 435-441, 2007) que la taille des grains de sel détermine la quantité de sel qui adhère aux frites. La forme, également, intervient : à masse de sel égale, ce sont des cristaux dont la surface est faible par rapport au volume, cubiques par exemple, qui adhèrent le mieux ; les cristaux en forme de plaquettes adhèrent moins, car ils se logent plus difficilement dans les anfractuosités microscopiques de la surface des objets frits.

Où trouver de petits cristaux de sel ? Il n'est ni difficile ni coûteux de mettre du sel de table habituel dans un moulin afin de produire du « sel glace » aux minuscules cristaux (on peut aussi utiliser le mortier et le pilon). Le sel ainsi produit tiendra bien aux frites... et sera plus facile à doser.

Le maintien de la température de friture, l'épongement des frites et l'utilisation de sel fin diminuent les dangers des fritures.



L'Environnement de Simulation Multiphysique

COMSOL Multiphysics constitue une solution complète à vos défis technologiques d'aujourd'hui : modélisation interactive, couplages multiphysiques illimités, import CAO.

Pour en savoir plus!

> Demandez votre CD gratuit des
Conférences COMSOL 2007 sur
www.comsol.fr/conference2007/cd



- Refroidissement en Electronique
- Processus Thermiques
- Acoustique et Vibration
- Bio-ingénierie
- Electromagnétisme
- Interaction Fluide-Structure

- Guides d'Ondes et Antennes
- Microfluidique
- Mécanique des Polymères
- Chauffage Résistif et Inductif
- MEMS
- CFD

www.comsol.fr

COMSOL, COMSOL MULTIPHYSICS, COMSOL REACTION ENGINEERING, LTD ET REPAIR SONT DES
MARQUES ENREGISTRÉES DE COMSOL AB. COMSOL SCRIPT EST UNE MARQUE DÉPOSÉE DE COMSOL AB.

Quand les riches maigrissent...

... les pauvres meurent, affirment un proverbe chinois et certaines théories économiques. La querelle sur une forte imposition des riches qui ne bénéficierait pas aux pauvres est d'actualité.



Il y a exactement cent ans, Anatole France publiait « L'Île des Pingouins ». Rappelons que saint Maël ayant baptisé par erreur une colonie de pingouins, Notre Seigneur dut les changer en hommes, les faisant ainsi rentrer dans l'Église et dans l'Histoire. Lors des premiers États de Pingouinie, saint Maël proposa à ses ouailles d'établir un impôt proportionnel, afin de subvenir aux dépenses publiques et à l'entretien des religieux : « Celui qui a cent bœufs en donnera dix ; celui qui en a dix en donnera un ». Morio, l'un des Pingouins plus riches, répondit au nom de tous les notables qu'ils étaient prêts à se dépouiller de tout ce qu'ils possédaient pour leurs frères Pingouins, mais qu'il fallait considérer uniquement l'intérêt public. « Or l'intérêt public exige de ne pas beaucoup demander à ceux qui possèdent beaucoup, car les pauvres vivent du bien des riches : c'est pourquoi ce bien est sacré. À prendre aux riches, vous ne retirerez pas grand profit, car ils ne sont guère nombreux ; et vous vous priveriez, au contraire, de toutes ressources, en plongeant le pays dans la misère... En chargeant tout le monde également et légèrement, vous épargnerez les pauvres, car vous leur laisserez le bien des riches ». Son discours fut salué par une ovation qui durait encore lorsque Greatauk, la main sur le pommeau de l'épée, fit cette brève déclaration : « Étant noble, je ne contribuerai pas, car contribuer est ignoble. C'est à la canaille à payer. »

Le débat est toujours d'actualité, même si les arguments se dissimulent dorénavant sous un jargon technique. Greatauk n'invoque plus les privilèges de la naissance, mais ceux de la juridiction : il a transféré son domicile fiscal en Belgique, et le siège social de son entreprise aux Îles Caïman. Il ne lui reste plus qu'à délocaliser l'entreprise elle-même, menace qu'il brandit plus efficacement qu'une épée. Morio évoque le PIB de la nation pingouine, et assure que le taux de croissance serait négativement affecté par une hausse des impôts, et surtout par un relèvement des tranches supérieures.

Le bon peuple raisonne comme si le PIB était réparti également, si bien que la croissance bénéficierait à toute la population, et repart convaincu. Il n'en est rien : c'est le paradoxe de la moyenne. En pesant ensemble 1 000 alouettes et un cheval, on trouvera que le poids moyen d'un animal est de un kilogramme, mais l'animal en question ne ressemble ni

aux alouettes ni au cheval. Si ce poids moyen augmente, c'est-à-dire si l'animal moyen prospère, il y a fort à parier que c'est le cheval qui engraisse. De même, une augmentation du PIB est parfaitement compatible avec un appauvrissement de la majorité et une concentration de la richesse entre les mains de quelques-uns, comme cela se passe en ce moment dans les pays anglo-saxons, où la part du dernier décile (les 10 pour cent les plus riches) dans les revenus totaux (hors plus-values) est passée de 32 pour cent en 1980 à 43 pour cent aujourd'hui, et la part du dernier centile de 8 pour cent à 17 pour cent.

La théorie économique a largement contribué à faire passer l'argent des riches pour l'argent des pauvres. Le problème central de cette théorie, tel qu'il a été formulé par Léon Walras vers la fin du XIX^e siècle et résolu par Kenneth Arrow et Gérard Debreu vers le milieu du XX^e, consiste à montrer qu'étant donnée une répartition initiale des richesses, il existe un système de prix d'équilibre. Ces prix déterminent la production et les échanges, l'offre est égale à la demande sur tous les marchés, et une fois le processus terminé, l'économie se trouve dans un nouvel état. Le grand mérite de cette nouvelle répartition des biens est qu'elle est préférable à la précédente (sans quoi les gens n'échangeraient pas) et qu'elle est efficace, au sens où on ne peut pas répartir de nouveau de manière à améliorer le sort de chacun.

Le miracle des prix d'équilibre est qu'ils permettent de donner aux uns sans retirer aux autres, mais une fois les échanges effectués, le miracle est terminé, et pour habiller Pierre il faut déshabiller Paul. L'impôt, disent les nantis, perturbe le fonctionnement normal du marché : les prix TTC ne conduisent plus à une situation efficace. Il y a donc du gaspillage : en pilotant l'économie d'une autre manière, on pourrait améliorer le sort de chacun sans nuire à personne.

On oublie trop souvent que les prix d'équilibre et l'état efficace auquel ils conduisent dépendent de la répartition des richesses. La concentration de la richesse entre les mains de quelques-uns ou l'équidistribution au sein de la population conduisent à des prix d'équilibre fort différents, et à des répartitions après échanges qui ne le seront pas moins, tout en étant efficaces l'une et l'autre. La conclusion est simple : le seul impôt irréprochable, en économie néoclassique, le seul qui ne perturbe pas le jeu du marché, c'est l'impôt sur les successions. Pourquoi donc est-ce le plus attaqué ?

JUSQU'OU IRONT LES DISPARITÉS DE REVENUS ET DE RICHESSES ?

« Un appel urgent
pour un réexamen
profond du capitalisme »
Joseph Stiglitz

SUPERCAPITALISME

Le choc entre le système
économique émergent
et la démocratie

Robert Reich

288 pages environ
978-2-7117-4338-4
Prix non communiqué
Éditions Vuibert

En librairie
le 17 Janvier 2008



« Le coupable, c'est le supercapitalisme — le système économique émergent où les consommateurs et les investisseurs détiennent pratiquement la totalité du pouvoir, n'en laissant qu'une partie congrue aux salariés et aux citoyens ». Un ouvrage majeur par l'économiste américain Robert Reich, ancien ministre de Bill Clinton.

www.vuibert.fr

Robert Reich
sera à Paris
le jeudi 17 janvier 2008,
où il donnera
une conférence
à Sciences-Po à 18h00

Inscription
auprès du MPA,
teresa.cal@sciences-po.fr

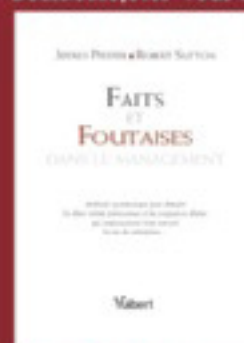
ET AUSSI...

Êtes-vous un sale con ?



208 p., 18 €, 7-4333
Vendu à plus de 25 000 ex.
entre avril et décembre 2007

Décideurs, êtes-vous formatés ?



288 p., 25 €, 7-4335

De l'insuline à l'insuline

Dans l'article *Le diabète du bébé* (voir *Pour la Science* n° 363, janvier 2008), un détail sur le fonctionnement des cellules bêta du pancréas me perturbe. Sur l'ensemble des schémas des cellules bêta, il est mentionné que des récepteurs à insuline permettent l'ouverture de canaux à glucose, et que l'entrée de ce glucose permet, après plusieurs étapes, l'exocytose d'insuline. Comment un tel système peut-il fonctionner ? De plus, sur la figure 2, dans le cas du diabète permanent, l'absence de libération d'insuline ne devrait-elle pas empêcher l'entrée de glucose dans les cellules bêta ?

Gilles Rocu, Aulnay-sous-Bois



Réponse de Michel Polak et Philippe Froguel

En fait, l'insuline augmente l'activité des transporteurs du glucose dans la plupart des tissus, mais pas dans tous. En l'absence d'insuline, la cellule bêta à insuline peut réagir à une gamme croissante de concentrations de glucose – et le laisser entrer –, car le transporteur du glucose y est différent de celui de la plupart des organes, qui nécessite de l'insuline. Le signal glucose est transformé dans la cellule bêta en un signal métabolique (avec l'ATP) qui contrôle à son tour, via les canaux potassiques et les variations de potentiel de membrane, la libération d'insuline.

Le système fonctionne donc bien si la cellule bêta libère suffisamment d'insuline pour permettre l'entrée de glucose dans les autres cellules du corps. Quand la cellule bêta est altérée, l'entrée de glucose est toujours possible, mais la libération d'insuline impossible.

Baia est Baïes

J'ai lu avec intérêt l'article sur Baia la romaine (voir *Pour la Science* n° 363, janvier 2008). Une remarque toutefois ; Baia est le nom italien de ce site, mais il existe un nom usuel français : Baïes.

Jean-Luc Gauchon, Roche

Pourquoi le Chinook est-il chaud ?

Dans son article *La fin des Dames anglaises ?* (voir *Pour la Science* n° 363, janvier 2008, p. 90), Jean-Paul Delahaye laisse entendre que le Chinook devrait sa chaleur à la compression due à la perte d'altitude. Cette explication est trop partielle ; si c'était le cas, tous les vents de montagne seraient chauds, alors que lors du franchissement d'un relief, la compression adiabatique – sans échange de chaleur – à la descente n'est que la symétrique de la détente adiabatique à la montée – ce qui explique pourquoi il fait plus froid en altitude. La spécificité des vents tels que le Chinook, et plus généralement des fœhns, est de gagner de la chaleur via la libération de chaleur latente de condensation lors de la montée (c'est pourquoi il pleut sur le versant au vent), ce qui n'a pas de symétrique à la descente, d'où la « fabrication » d'un vent chaud (et sec).

Fabrice Neyret, par courriel

Des cafards robotisés

En ce qui concerne l'actualité scientifique *Des cafards dirigés par des robots* (voir *Pour la Science* n° 363, janvier 2008, p. 23), je ne comprends pas l'engouement que ce travail suscite. Certes, il est suffisamment bien construit pour qu'on s'y laisse prendre. Mais il suffit de quelques minutes de réflexion pour se rendre compte que la clé de toute l'expérimentation réside dans le « parfum de cafard », sous forme de phéromones, qui est déposé sur les robots. Imaginez simplement qu'il n'y ait pas eu de robot et qu'on ait déposé la goutte sous l'abri le plus éclairé. Nul doute que les cafards s'y seraient tous retrouvés.

Richard-Emmanuel Eastes, École normale supérieure, Paris

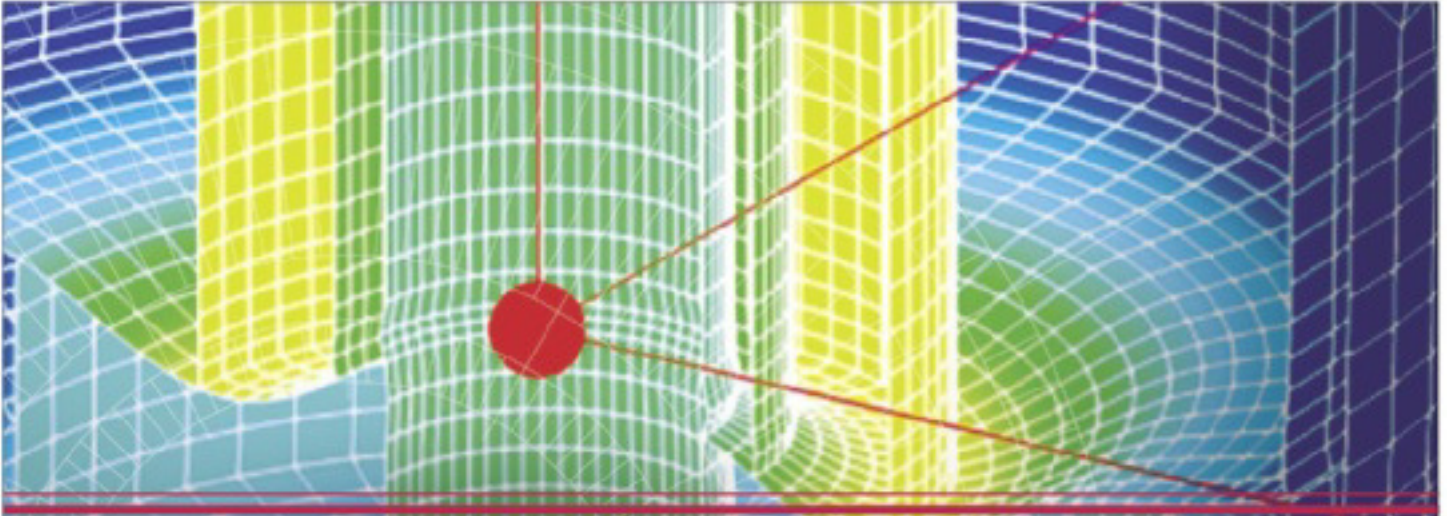
Réponse de José Halloy, biologiste à l'Université libre de Bruxelles

En effet, une goutte de l'odeur appropriée de cafard sous un abri suffit à favoriser l'agrégation des cafards à cet endroit. C'est d'ailleurs ce test expérimental qui nous a permis d'identifier le mélange de molécules qui attire les cafards ! Ce mélange est extrait des nombreuses molécules que le cafard porte sur son squelette externe. Dans notre étude, la différence avec cette situation est que les robots sont mobiles : il s'agit d'un système dynamique. Contrairement au dépôt d'une phéromone à un endroit spécifique, notre situation est plus proche de la situation naturelle où les cafards ne déposent pas de gouttes de phéromones, mais « transpirent » ces odeurs. L'abri sous lequel les groupes s'établissent n'est pas prédéfini, mais émerge des interactions entre les individus.

Dans les cas de deux abris d'obscurité équivalente, les cafards se regroupent tous dans l'un ou l'autre des abris, au hasard ; il n'y a pas de raison de préférer l'un plutôt que l'autre. Si nous introduisons des robots programmés pour se comporter comme les cafards, nous observons la même chose : le groupe mixte (des cafards et des robots) choisit dans 50 pour cent des cas un abri, dans 50 pour cent des cas, l'autre. La sélection de l'abri émerge des échanges entre les individus naturels et artificiels. Dans cette expérience, nous ne pouvons pas prédire quel abri sera sélectionné ; il existe alors deux états stables pour ce système dynamique.

Dans le cas d'un abri clair et d'un abri sombre, en condition naturelle, les cafards ont tendance à se regrouper dans l'abri sombre. Dans nos expériences, ils le font dans 73 pour cent des cas, contre 27 pour cent des cas dans l'abri clair. Si nous introduisons des robots programmés pour préférer l'abri clair, la solution la plus fréquente change : le groupe mixte se regroupe dans 61 pour cent des cas dans l'abri clair, contre 27 pour cent en l'absence des robots. Cependant, dans 39 pour cent des cas, les robots se retrouvent avec les blattes dans l'abri foncé, et ce malgré leur préférence individuelle pour l'abri clair. En effet, la sélection des abris dépend non seulement des préférences individuelles, mais aussi des interactions sociales. Nous avons donc un système dynamique, en accord avec un modèle mathématique, produisant des solutions nouvelles qui émergent des interactions entre les individus naturels et artificiels. Voilà une situation plus intéressante que le simple dépôt de gouttes de phéromones à certains endroits !

Pour nous écrire : tribune@pouirlascience.fr



Des formations à la pointe du savoir

IPSI

Nos sessions 2008

**Toutes nos sessions se déroulent
à la Cité des sciences et de l'industrie à Paris**

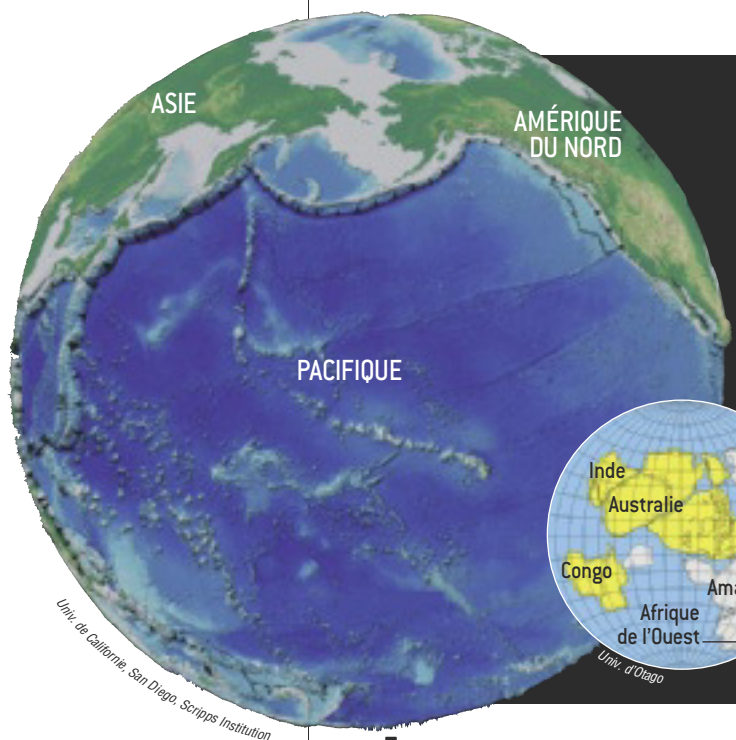
- ▶ Optimisation multidisciplinaire
- ▶ Veille technologique / Intelligence économique
- ▶ Simulation multiphysique
- ▶ Gestion de l'innovation
- ▶ Tribologie / Surface fonctionnelle
- ▶ Prospectives marketing et méthodes
- ▶ Traitement du signal d'image
- ▶ Biomécanique / Biomédical
- ▶ Conception sensorielle
- ▶ Fatigue de contact

Notre catalogue formation est disponible sur simple demande

CONTACT

Renseignements techniques : **Laurence Honnorat**
Inscriptions : **Service formation** Tél. : 03 44 67 31 45
ipsi@cetim.fr - www.ipsi-cetim.fr





La croûte terrestre est constituée de plaques (ci-dessus) qui se déplacent au gré de l'expansion des bassins océaniques et des zones de subduction, où une plaque passe sous une autre. Dans 350 millions d'années, quand le Pacifique (à gauche), dont le pourtour est riche en zones de subduction, aura disparu, ces mouvements cesseront probablement, comme ils l'ont fait il y a un milliard d'années, à l'époque du supercontinent Rodinia (ci-contre, ici peu après son éclatement).

La tectonique intermittente

En 1912, l'astronome et météorologue allemand Alfred Wegener exposait sa théorie de la dérive des continents, que l'on peut résumer ainsi dans sa version actuelle : la croûte terrestre est morcelée en plaques – océaniques ou continentales – qui « flottent » à la surface du manteau terrestre. Ces déplacements, de quelques centimètres par an, sont dus à un flux de chaleur interne évacuée surtout par convection, c'est-à-dire par la mise en mouvement des roches du manteau.

La plupart des géophysiciens pensaient le phénomène lent, certes, mais continu. Ce n'est pas le cas, à en croire Paul Silver, de l'Institut Carnegie de Washington, et Mark Behn, de l'Institut océanographique Woods Hole.

La croûte terrestre naît au niveau de dorsales océaniques où des roches du manteau remontent. Ce faisant, des océans apparaissent et s'élargissent, ce qui est aujourd'hui le cas de l'Atlantique. À cause de cette expansion, des plaques se rapprochent et entrent en contact. Dans certains cas, elles se heurtent et donnent naissance à des chaînes de montagnes, par exemple l'Himalaya à la frontière des plaques indienne et eurasiennne. Dans d'autres cas, une plaque se glisse sous l'autre, moins dense, et rejoint le manteau ; dans ces zones dites de subduction, le volcanisme est intense.

Ces phénomènes façonnent notre planète. Ainsi, il y a 200 millions d'années, toutes les terres émer-

gées étaient réunies en un seul supercontinent, la Pangée. Les modèles prédisent que dans 350 millions d'années le Pacifique disparaîtra. Or le bassin de cet océan regroupe les principales zones de subduction connues aujourd'hui : en leur absence, la croûte terrestre s'immobiliserait jusqu'à ce que de nouvelles zones de subduction se créent. Mais on ignore le mécanisme de leur apparition.

Ainsi, il y a 30 à 50 millions d'années, la collision de l'Afrique et de l'Inde avec l'Eurasie a rayé de la carte le bassin océanique de la Téthys et, depuis, aucune nouvelle zone de subduction n'est apparue. Si un tel calme apparent continue à régner, après la fermeture du Pacifique, les plaques seront « gelées ». En étudiant différents indicateurs géochimiques dans les roches des terres émergées (la distribution de l'hélium 3 et de l'hélium 4, les proportions de niobium et de thorium), les géophysiciens ont montré qu'une telle situation avait déjà eu lieu il y a environ un milliard d'années, quand un seul supercontinent, nommé Rodinia, existait.

Quelles sont les conséquences d'une tectonique intermittente ? D'abord, les épisodes d'inactivité ralentiraient le refroidissement de la Terre, ce qui expliquerait pourquoi les modèles peinent à reconstituer l'histoire thermique de notre planète. Ensuite, il faut peut-être revoir les scénarios d'évolution des continents.

Loïc Mangin

Science, vol. 319, pp. 85-88, 2007

Réparer les muscles des myopathes ?

En France, un garçon sur 3 500 est atteint de myopathie de Duchenne, une dégénérescence progressive des muscles due à une anomalie du gène DMD (situé sur le chromosome X) qui code la dystrophine. Cette protéine assure la tenue et la cohésion des fibres musculaires. Sans elle, le muscle ne peut résister aux forces exercées lors des contractions et finit par s'atrophier. Les enfants souffrent d'abord de faiblesse musculaire, sont contraints de se déplacer en fauteuil roulant dès huit ou neuf ans, puis les muscles respiratoires eux-mêmes sont atteints ; bien pris en charge, ils vivent jusqu'à 20 ou 30 ans. Aucun traitement n'est disponible et cette maladie est l'objet de nombreuses recherches. Ainsi, l'équipe de Luis Garcia, de l'Institut de myologie à Paris, et celle de Yvan Torrente, à l'Université de Milan, ont réussi à restaurer la production de dystrophine dans des cellules de malades en culture. Ils ont réinjecté ces cellules humaines modifiées à des souris malades, dont l'état s'est amélioré.

Le gène DMD codant la dystrophine est le plus long du génome. Il est transcrit en une molécule d'ARN qui comporte des régions dites codantes (ou exons), mises bout à bout lors du phénomène d'épissage pour donner une molécule d'ARN messager traduite en dystrophine. Dans la myopathie, des mutations du gène

entraînent des défauts dans la lecture de l'ARN messager : la dystrophine n'est plus synthétisée.

Comment les chercheurs ont-ils rétabli la production de dystrophine ? Ils ont prélevé dans le sang ou dans les muscles de malades atteints de myopathie de Duchenne des cellules souches musculaires – des cellules immatures prêtes à se transformer en fibres musculaires et qui contiennent donc les anomalies du gène DMD. Ils ont corrigé en culture ces anomalies par la technique du « saut d'exon » : au moment de l'épissage, les exons mutés sont éliminés et, par conséquent, sont absents de l'ARN messager. Pour ce faire, un virus (rendu inoffensif) introduit dans la cellule souche une molécule qui masque les exons mutés. La machinerie qui fabrique les protéines ignore les exons mutés et « saute » par-dessus les anomalies ; une dystrophine tronquée, mais fonctionnelle, est alors fabriquée. Ces cellules souches ont ensuite été injectées dans des souris modèles de la myopathie de Duchenne : 45 jours après, les cellules des souris expriment la dystrophine et les animaux se déplacent mieux. Si l'on arrive à transposer la technique à l'homme, les patients pourront alors refabriquer l'indispensable protéine à partir de leurs propres cellules souches musculaires réparées.

Bénédicte Salthun-Lassalle

Cell Stem Cell, vol. 1, pp. 646-657, 2007

17 milliards de soleils

Telle est la masse colossale du plus gros trou noir jamais repéré ! Équivalente à celle d'une petite galaxie, elle vient d'être mesurée grâce au mouvement d'un trou noir plus petit – 100 millions de masses solaires tout de même – qui gravite autour. Le monstre réside dans une galaxie située à 3,5 milliards d'années-lumière.

Attention soutenue

Quelque chose surgit dans votre champ visuel et attire votre attention, alors que vous étiez concentré sur autre chose. Peut-on être attentif à plusieurs phénomènes en même temps ? Des chercheurs français et américains viennent de montrer qu'il n'existe qu'un seul faisceau attentionnel, mais il se déplace sept fois par seconde... de quoi disperser l'attention !

Climats tempérés sur planètes excentriques

L'ellipticité de l'orbite de certaines exoplanètes favoriserait l'existence de zones tempérées à leur surface.

L Le climat est-il tempéré sur les exoplanètes qui présentent toujours la même face à leur étoile ? *A priori*, non. Toutefois, des simulations menées par Anthony Dobrovolskis, du Centre de recherche AMES de la NASA, en Californie, suggèrent que sur celles dont l'orbite est suffisamment elliptique, il existe des régions où l'ensoleillement est modéré.

En raison des interactions gravitationnelles, les planètes rocheuses proches de leur étoile – comme c'est le cas pour nombre d'exoplanètes – sont souvent en résonance synchrone, c'est-à-dire que leur période de révolution autour de l'étoile et celle de rotation propre sont égales. Dans cette configuration, la planète présente toujours la même face à son étoile, comme la Lune par rapport à la Terre. Les exoplanètes auraient ainsi un hémisphère toujours ensoleillé, où règne la fournaise, et un autre perpétuellement plongé dans l'obscurité et glacé...

Cet enfer serait cependant plus tempéré qu'il n'y paraît. Lorsque l'orbite est elliptique, la vitesse de la planète varie. Au plus près de l'étoile, sa vitesse angulaire de révolution est plus élevée que la rotation propre ; vue de l'étoile, la planète paraît alors avoir légèrement tourné sur elle-même. Au point le plus éloigné, elle semble avoir tourné dans l'autre sens.

Ce mouvement dit de libration décale la face exposée à l'étoile de quelques degrés dans un sens puis dans l'autre. Il est d'autant plus ample que l'orbite est elliptique – ou « excentrique ». A. Dobrovolskis a calculé que la libration atteint 60 degrés pour une excentricité de 0,5, une valeur commune pour les exoplanètes (en comparaison, celle de la Terre est de 0,0167).

Il a ensuite montré que cette bascule fait apparaître à la surface de la planète une large bande tantôt ensoleillée, tantôt dans l'obs-

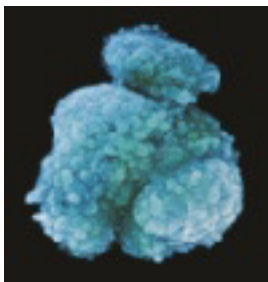
curité, à la frontière entre le côté jour et le côté nuit. Dans cette région, la température serait clémente. Il a également montré que les profils d'ensoleillement correspondant à d'autres résonances possibles (non synchrones) protègent aussi certaines longitudes des températures extrêmes. Suffisamment pour abriter la vie ?

Philippe Ribeau-Gésippe

Icarus, vol. 192 (1), pp. 1-23, 2007

Même si une planète présente toujours la même face à son étoile, l'ellipticité de son orbite fait apparaître des bandes tempérées le long de la frontière entre le côté jour et le côté nuit.





© Biophoto associates photo researchers, Inc.

Le chromosome Y (en bleu) est atrophié par rapport au chromosome X (en rose). Les cellules femelles, dotées de deux X, en inactivent un. Une région chromosomique impliquée dans cette inactivation a été identifiée.

Parité chromosomique

Au moment de la fécondation, le spermatozoïde apporte un chromosome sexuel X ou Y à l'ovule qui contient déjà un X. Le sexe du futur individu est à ce moment décidé : XY donnera un garçon, XX une fille. Cependant, le chromosome Y est beaucoup plus petit que le X et contient donc moins de gènes : les individus XX ont près de 2 000 gènes de plus que les XY. Ce déséquilibre est compensé par l'inactivation d'un chromosome X dans l'embryon femelle, un phénomène mal connu. Edith Heard et ses collègues (Unité CNRS/Institut Curie, Dynamique nucléaire et plasticité du génome), ont découvert comment les cellules « éteignent » un de leurs chromosomes X dès qu'elles constatent qu'elles en ont deux.

Les biologistes ont identifié une région, nommée *Xpr*, sur le chromosome X. Dans une cellule femelle, les deux chromo-

somes X s'accrochent temporairement par cette région. Par ce contact, la cellule détecte qu'elle contient deux X et le processus d'inactivation peut alors se déclencher. De quelle façon ?

Xpr appartient à une région connue pour gérer les étapes ultérieures du processus d'inactivation : le centre d'inactivation du X. Ce dernier est notamment capable de produire deux ARN non codants, l'ARN *Xist* et son antisens, l'ARN *Tsix/Xite*, dont l'expression est initialement équivalente sur les deux X. Dès que *Xpr* repère la présence de plusieurs chromosomes X, *Xists* s'accumule sur l'un des deux chromosomes X, choisi au hasard, et entraîne, par un mécanisme inconnu, son inactivation. Le X inactif adopte alors une forme condensée et les quelque 2 000 gènes qu'il contient sont réduits au silence.

Le choix du X inactivé est aléatoire et a lieu lorsque l'em-

bryon est constitué de quelques cellules. Ainsi, un organisme femelle se compose de cellules où le X paternel s'exprime et d'autres où c'est le X maternel.

Le rôle de la région *Xpr* dans l'inactivation du X a été confirmé par l'introduction d'un petit fragment d'ADN le contenant dans une cellule mâle : la cellule XY leurrée croit recenser plusieurs X et inactive son unique X.

Ce fonctionnement n'est pas nécessairement propre aux chromosomes sexuels. En novembre 2007, une étude a montré que pour dix pour cent des gènes (notamment plus que ce que l'on pensait), seule une des deux versions (celle du père ou celle de la mère) est transcrite. Cette sélection aléatoire est peut-être elle aussi fondée sur des mécanismes de contacts chromosomiques similaires à ceux utilisés par *Xpr*.

L. M.

Science, vol. 318, pp. 1632-1636, 2007

Du gaz plein les comètes

L'analyse par une équipe française des échantillons de la comète *Wild2*, ramenés par la mission *Stardust*, le montre : les comètes sont riches en gaz rares. Cela renforce l'idée que les comètes ont apporté des éléments volatils sur Terre. Par ailleurs, certains composés ne se forment qu'à haute température, ce qui suggère que les comètes sont nées près du Soleil avant d'être repoussées au loin.

Pas de miel sur les plaies

Antiseptique, le miel ? En tout cas, il n'améliore pas les pansements, concluent des chirurgiens néozélandais après avoir bandé, avec des pansements imprégnés de miel, plusieurs centaines de patients souffrant d'un ulcère de la jambe. Les pansements au miel ont entraîné nettement plus de complications que les pansements sans miel du groupe de référence.

L'ancêtre terrestre des baleines

Depuis Darwin, les paléontologues savent que les cétacés, telles les baleines, ont pour ancêtres un artiodactyle, c'est-à-dire un ongulé à nombre pair de doigts (hippopotame, chameau...). Cependant, on ignorait à quoi ressemblait leur aïeul terrestre. Le fossile étudié par Hans Thewissen, de l'Université du Nord-Est de l'Ohio, comble cette lacune. En effet, *Indohyus*, un artiodactyle fossile de 48 millions d'années et mis au jour en Inde, partage plusieurs caractéristiques, notamment crâniennes et auriculaires, avec les baleines. L'animal, une sorte de daim de la taille d'un renard, aurait consommé des végétaux terrestres, mais il aurait passé le plus clair de son temps dans l'eau, probablement pour se mettre à l'abri des prédateurs.

L. M.

Nature, vol. 450, pp. 1190-1194, 2007



© C. Buell

La Lune rajeunie de 20 millions d'années

De nouvelles datations de roches lunaires réalisées par Mathieu Touboul, de l'Institut fédéral de technologie de Zurich, suggèrent que notre satellite, né de l'impact d'un planétésimal avec la Terre, est 20 millions d'années plus jeune que ce que l'on pensait.

L'impact était supposé avoir eu lieu 30 millions d'années à peine après la naissance du Système solaire, il y a 4,567 milliards d'années ; la Lune s'est ensuite solidifiée en dix millions d'années. Cette datation s'appuie notamment sur un radio-isotope, le tungstène 182, produit par la désintégration de l'hafnium 182 et du

tantalum 182. Or ce dernier n'était pas initialement présent : il résulte de l'impact des rayons cosmiques sur la surface lunaire. Il ne doit donc pas être pris en compte dans la datation. La datation corrigée suggère que l'impact a eu lieu 50 millions d'années après la naissance du Système solaire. Cette nouvelle chronologie recule d'autant la fin de la formation des planètes telluriques par accrétion, et remet en question l'idée selon laquelle le noyau terrestre s'est formé suite à l'impact géant.

Ph. R.-G.

Nature, vol. 450, pp. 1206-1209, 2007

Esprit coopératif

John McNamara, de l'Université de Bristol, a créé un modèle pour expliquer le succès évolutif des individus coopératifs. *A priori*, les lignées qui les produisent devraient s'éteindre, car leur succès reproductif serait obéré par leur tendance à se faire exploiter. Dans le modèle, cependant, l'attitude plus ou moins sélective d'un individu dans sa recherche de partenaires détermine le niveau de coopération qu'il attend. L'esprit coopératif faisant les bons partenaires, même une population initialement dépourvue d'individus coopératifs évolue, d'après le modèle, vers une population comptant des individus coopératifs.

Effet Casimir dans un liquide

Deux plaques métalliques parallèles placées dans le vide et espacées de quelques micromètres s'attirent légèrement : c'est l'effet Casimir, qui est dû aux fluctuations quantiques du champ électromagnétique – des fluctuations toujours présentes, même au zéro absolu et en l'absence de charges électriques. Christopher Hertlein et ses collègues, à l'Université de Stuttgart, viennent de mesurer directement une force analogue, mais due aux fluctuations des concentrations locales dans un mélange liquide.

Cette « force de Casimir critique » avait été prédite en 1978 par le Britannique Michael Fisher et le Français Pierre-Gilles de Gennes, et l'on avait depuis certaines indications de son existence. L'adjectif « critique » se rapporte aux systèmes qui changent d'état de façon continue lorsque certains paramètres, la température par exemple, atteignent des valeurs dites critiques, le phénomène s'accompagnant de grandes fluctuations (par exemple de la densité dans une transition liquide-gaz). Le système choisi par les physiciens allemands est un mélange d'eau et de 2,6-lutidine, un liquide organique. Lorsque la concentration en masse de lutidine est égale à 28,6 pour cent, le mélange se dissocie au-dessus d'environ 34 °C, point critique près duquel les fluctuations locales de concentration deviennent très importantes.

L'équipe de Stuttgart a mesuré, par une méthode de microscopie optique, la force s'exerçant entre une surface de silice et une microsphère de polystyrène, séparées par environ 0,1 micromètre de mélange liquide. Cette force est attractive ou répulsive selon l'affinité des surfaces pour l'eau ou la lutidine. Les résultats sont en très bon accord avec les estimations théoriques et constituent une prouesse : l'intensité des forces ne dépasse pas une fraction de piconewton (10^{-12} newton), mais a pourtant été mesurée au femtonewton (10^{-15} newton) près.

Maurice Mashaal

Nature, vol. 451, pp. 172-175, 2008

Des canaux membranaires à ouverture « électronique »

Une équipe britannique a découvert un troisième type de mécanisme d'ouverture des canaux ioniques.

Nos cellules sont pourvues de multiples canaux ioniques, c'est-à-dire des protéines qui constituent des portes entre deux compartiments (le plus souvent les milieux intra- et extracellulaires) par lesquelles passent des ions. Selon la protéine, le canal n'est perméable qu'à un seul type d'ions ou à plusieurs. Ces mouvements d'ions sont essentiels pour la transmission de l'influx nerveux, les contractions musculaires, le maintien de la glycémie, etc.

Un canal ionique peut être ouvert ou fermé, et l'on distinguait jusqu'à aujourd'hui deux principaux types de mécanismes d'ouverture. Celle des canaux voltage-dépendants résulte de la variation de la polarité membranaire (la répartition de charges électriques de part et d'autre de la membrane où est enchâssé le canal), alors que les canaux ligand-dépendants s'ouvrent quand leur molécule « attirée » se lie à eux. David Beech et ses collègues de l'Université de Leeds, au Royaume-Uni, ont mis en évidence un nouveau mécanisme d'ouverture : un canal ionique s'ouvre lorsqu'une protéine extracellulaire lui transfère un électron.

Les biologistes se sont intéressés aux canaux TRP, des tunnels à cations (des ions positifs tels le sodium et le calcium), codés chez l'humain par 28 gènes distincts. L'un de ces canaux TRP (TRPV1) est responsable de la perception de chaud quand on mange un piment, un autre (TRPM8), de celle du froid et de la sensation due au menthol. Ici, ce sont les canaux TRPC5 et TRPC5-TRPC1 qui étaient au centre des investigations.

Plusieurs expériences ont révélé que ces canaux s'ouvrent grâce à un mécanisme original mettant en jeu la thiorédoxine. Dans le milieu extracellulaire, cette protéine fournit des électrons au canal (le canal subit une réduction, tandis que la protéine est oxydée). Ce faisant, la thiorédoxine coupe un pont disulfure (deux atomes de soufre liés) et modifie la configuration de la protéine-canal, qui passe d'un état fermé à l'état ouvert : les ions passent et traversent donc des canaux d'un nouveau type, des canaux dont l'ouverture est commandée par une oxydo-réduction.

D'où vient la thiorédoxine ? Elle est normalement fabriquée dans la cellule où elle piège les radicaux libres, ce qui protège la cellule de leurs effets délétères. Cependant, elle est parfois sécrétée dans le milieu extracellulaire où elle participe à des processus anti-inflammatoires. En particulier, elle y est abondante chez les malades atteints de polyarthrite rhumatoïde (une inflammation chronique des articulations). En ouvrant les canaux TRPC5 et TRPC5-TRPC1, la thiorédoxine extracellulaire empêcherait la libération d'enzymes qui, en remodelant les tissus articulaires, favorisent la progression de la maladie. L'ubiquité de la thiorédoxine et des canaux TRP dans l'organisme est peut-être le signe que ces canaux d'un nouveau type ne sont pas restreints aux conditions pathologiques associées à la polyarthrite rhumatoïde.

L. M.

Nature, vol. 451, pp. 69-72, 2008

Le trésor du « bout du monde »

Une fouille d'archéologie préventive a mis au jour le plus grand trésor gaulois jamais découvert en Bretagne : un lot de 545 pièces de monnaie.

A la fin du III^e siècle avant notre ère, un notable s'installe à Rosquelfen, non loin de Laniscat, dans les Côtes d'Armor. Au I^{er} siècle avant notre ère, l'un de ses successeurs sur place y cache 545 pièces (statères). Découvert fin 2007 par l'équipe d'Eddie Roy, de l'INRAP, ce dépôt monétaire est le plus grand trésor gaulois jamais trouvé en Bretagne.

Le trésor se compose de 58 statères et de 487 quarts de statères, tous en électrum, un alliage d'or, d'argent et de cuivre. Il fut caché à côté de l'une des dépendances de la demeure, peut-être dans l'une des céramiques dont les éclats ont été retrouvés parmi les pièces. Son importance illustre la richesse exceptionnelle à l'époque gauloise que pouvait détenir un aristocrate de cette région.

Denis Gliksmann / INRAP



L'établissement agricole de Rosquelfen – indéniablement aristocratique avec son vaste enclos de 7 500 mètres carrés – se trouve sur le territoire des Osismes. Ces Gaulois, dont le nom signifierait « ceux du bout du monde » vivaient à l'extrémité de la péninsule bretonne. Au I^{er} siècle avant notre ère, leur chef-lieu était Carhaix dans le Finistère, qu'une voie romaine, succédant sans doute à une voie gauloise plus ancienne, reliait à Rennes. On a longtemps pensé que leurs voisins les Vénètes, qui ont opposé une très forte résistance à César dans le Morbihan durant la guerre des Gaules, dominaient leurs voisins osismes. Toutefois, nombre de données archéologiques récentes suggèrent au contraire la prédominance des Osismes, qui

maîtrisaient le trafic maritime entre l'Atlantique et la Manche et possédaient des gisements de métaux précieux.

Or le trésor trouvé à Rosquelfen est composé de statères, type de monnaie qui servait plus à récompenser les guerriers qu'aux échanges commerciaux (ces statères ont en effet été introduits en Gaule par des mercenaires revenus de chez les Grecs à l'époque de Philippe II de Macédoine). La datation du trésor le fait remonter aux années de la conquête romaine, époque à laquelle l'imposante ferme gauloise fut remplacée par un enclos gallo-romain. De là à penser que le trésor a été préparé par un chef pour des guerriers osismes qu'il emmenait se battre contre les Romains, il n'y a qu'un pas... que les archéologues ne franchiront que si l'étude numismatique approfondie du trésor le confirme.

François Savatier



Évacuer une foule au plus vite

Où doivent être placées les issues de secours ? De quelle largeur doivent-elles être ? De telles questions tourmentent les architectes et concepteurs de lieux publics, mais aussi, depuis plusieurs années, les physiciens qui y trouvent un terrain d'application de leurs concepts et méthodes.

Daichi Yanagisawa et Katsuhiro Nishinari, de la Faculté d'ingénierie de l'Université de Tokyo, viennent d'élaborer un modèle de déplacement de piétons près d'une sortie ; il leur a permis d'exprimer mathématiquement, pour la première fois, le flux moyen de piétons évacués en fonction de la largeur de l'issue, de sa position le long du mur et du caractère, coopératif ou concurrentiel, des piétons.

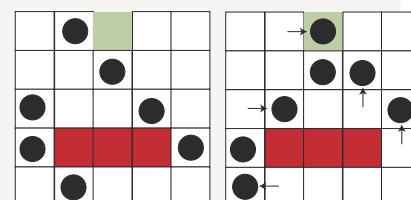
Ce modèle est un perfectionnement de modèles précédents, de type automates cellulaires : les piétons se déplacent sur un damier en avançant d'une case au plus à la fois dans l'une des

quatre directions (avant, arrière, gauche, droite), chaque case ne supportant qu'une seule personne (*voir la figure*). Tout déplacement est affecté d'une probabilité, dont la valeur dépend notamment de la distance entre la case convoitée et la sortie. Un paramètre de frottement μ , compris entre 0 et 1, modélise les situations de conflit, quand plusieurs piétons visent la même case : avec la probabilité μ , tous les piétons impliqués restent sur place, tandis qu'avec la probabilité $1 - \mu$, le déplacement de l'un d'eux, choisi au hasard, est autorisé. Ce paramètre traduit le degré de coopération (μ petit) ou de concurrence (μ grand) entre personnes.

Les deux chercheurs japonais ont introduit un nouveau paramètre, inspiré de l'observation que les gens marchent d'autant plus vite qu'ils se situent plus loin de la sortie ; ce paramètre contrôle la vitesse de déplacement près de l'issue. Ils ont alors pu expri-

mer le flux moyen de piétons évacués en fonction des divers paramètres, et ces calculs sont en accord avec les simulations. Le modèle montre notamment qu'en situation de concurrence, il est préférable de placer l'issue de secours au coin d'une pièce plutôt qu'au milieu d'un mur : un résultat utile et peu intuitif. M. M.

Phys. Rev. E, vol. 76, article 061117, 2007



Un exemple de deux états successifs dans un modèle de piétons se dirigeant vers la sortie (en vert), en présence d'obstacles (en rouge). Les spécificités d'un tel modèle d'automates cellulaires résident dans la façon dont sont choisies les probabilités de chaque déplacement possible.

Des singes doués en calcul

Les gènes, les outils, la culture... Un à un, les candidats au statut de « propre de l'homme » tombent et, à chaque fois, la frontière entre humains et primates s'estompe. Dernier bastion à terre, le calcul mental : Jessica Cantlon et Elizabeth Brannon, de l'Université Duke, aux États-Unis, ont montré que des macaques rhésus sont presque aussi performants que des étudiants pour effectuer de petites additions.

Face à un écran, deux primates, Feinstein et Boxer (du nom de deux sénateurs américains), ainsi que des étudiants, regardaient successivement deux ensembles de points avant de se voir proposer deux cases contenant des points. Ils étaient récompensés lorsqu'ils désignaient celle dont le nombre de points correspondait à la somme des deux ensembles précédents.

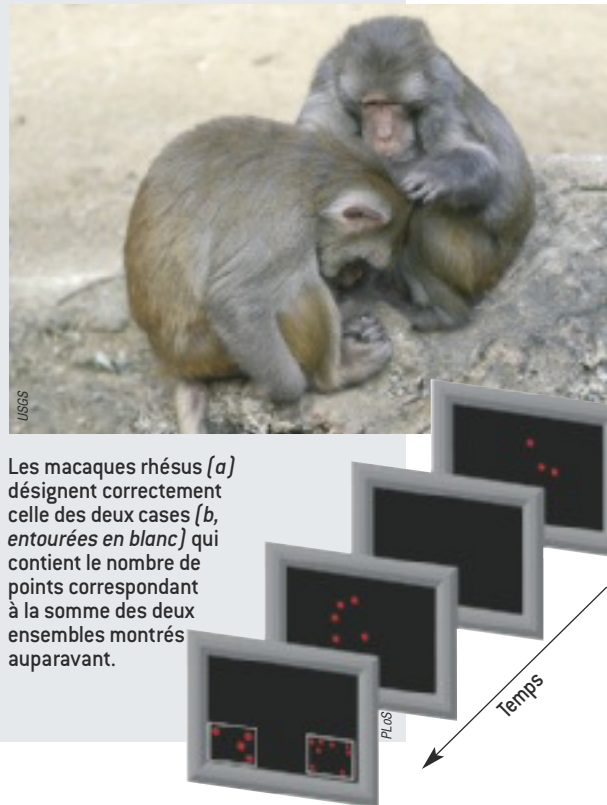
Les nombres utilisés étaient inférieurs à 20. Les étudiants ont répondu

correctement dans 94 pour cent des cas, contre 76 chez les macaques. Tous répondaient en une seconde, et avaient plus de difficultés quand les sommes correspondaient à des grands nombres, ou bien quand les deux propositions différaient peu. Ainsi, les singes peuvent effectuer des additions mentalement. Ces résultats confirment un peu plus encore que les animaux et les humains ont en commun diverses facultés arithmétique : par exemple, des études ont montré que plusieurs espèces peuvent comparer deux groupes d'objets.

Finalement, le propre de l'homme est peut-être à chercher dans les récompenses proposées par les expérimentateurs. Quand l'individu testé désignait la bonne case, il était récompensé, selon son espèce, par une boisson sucrée ou par dix dollars !

L. M.

PloS Biology, vol. 5, pp. 2912-2919, 2007



Les macaques rhésus (a) désignent correctement celle des deux cases (b, entourées en blanc) qui contient le nombre de points correspondant à la somme des deux ensembles montrés auparavant.

Des spermatozoïdes invisibles

Si un ovule et un spermatozoïde peuvent fusionner pour donner un œuf, c'est parce que le corps de la femme laisse la possibilité au spermatozoïde de rencontrer l'ovule. Ce n'est pas si évident, car tout corps étranger est normalement détruit par le système de défense immunitaire de la femme. Comment les spermatozoïdes peuvent-ils alors rester dans l'utérus plus de 72 heures ?

Poh-Choo Pang, de l'Imperial College à Londres, et ses collègues ont identifié les molécules – des sucres complexes nommés N-glycanes –, présentes à la surface des spermatozoïdes, qui les rendent invisibles aux cellules immunitaires chargées de détruire les corps étrangers. Ces molécules sont universelles, c'est-à-dire qu'elles trompent tous les systèmes immunitaires. En outre, elles se retrouvent sur certaines cellules cancéreuses, certains parasites et sur les cellules infectées par le virus du SIDA. Voilà pourquoi ces intrus passeraient inaperçus. Reste à comprendre comment ces molécules leurrent le système immunitaire.

B. S.-L.

The J. of Biol. Chem., vol. 282, pp. 36593-36602, 2007

Crotale, parfum pour écureuils

Comment éviter de devenir la proie d'un serpent ? La question est cruciale pour les écureuils terrestres. Aussi ont-ils recours à différentes stratégies dont une vient d'être découverte par Barbara Clucas, de l'Université de Californie, à Davis.

En étudiant deux espèces de rongeurs (*Spermophilus beecheyi* et *S. variagatus*), B. Clucas a montré que ces animaux mâchent les mues de crotales, puis se lèchent le pelage de façon à l'ordre de fragrances reptiliennes. Des observations ont révélé que les jeunes et les femelles, plus vulnérables, sont plus enclins à adopter cette stratégie. Les écureuils masquent ainsi leur propre odeur et restent dissimulés à l'odorat du prédateur particulièrement la nuit pendant leur sommeil. Un autre moyen de diversion, découvert il y a quelques mois, consiste à augmenter notablement la température de leur queue, ce qui perturbe les récepteurs à infrarouges des serpents. Enfin, les écureuils seraient insensibles aux venins de serpents.

L. M.

Animal Behaviour, vol. 75, pp. 299-307, 2007

Pluie ou pas pluie ?

Lundi au soleil et dimanche sous la pluie... Pleut-il davantage certains jours ? Non, d'après une étude finlandaise fondée sur des relevés météo. Mais une équipe américaine, à l'aide d'un satellite, révèle que l'été, dans le Sud-Est américain, la semaine est plus arrosée que le dimanche en raison des rejets industriels qui déclenchent la formation de gouttelettes. Qui croire ? Dans le doute, prenez votre parapluie.

Palmier géant

Avec ses feuilles de cinq mètres de diamètre et son tronc culminant à 18 mètres, il est même visible sur Google Earth : un palmier géant du genre *Chuniophoenix* vient d'être découvert à Madagascar. Il étonne non seulement par sa taille, mais aussi par une floraison durant laquelle il dépense toutes ses ressources. Découvert par hasard dans la campagne malgache, il est en cours de description.



Le Jupiter de Doliché...

...a enfin été découvert à Doliché,
site d'origine de son culte.

Doliché est une ville antique proche de l'Euphrate, non loin de la ville turque de Gaziantep. Siècle d'un très ancien culte au dieu mésopotamien Hadad, elle devint ensuite celui d'un culte au dieu syriaque Baal. Quand, dans le dernier tiers du I^{er} siècle de notre ère, les Romains conquièrent les lieux, ils superposèrent le culte de Jupiter à celui du dieu de Doliché. Curieusement, le culte de Doliché a ensuite été diffusé par les soldats dans tout l'empire sous le nom de *Jupiter Dolichenus*.

On le sait notamment parce qu'on a retrouvé des représentations du dieu de Doliché dans nombre d'endroits de l'empire, sauf... à Doliché. Cette lacune vient d'être comblée par une équipe germano-turque dirigée par Engelbert Winter, de l'Université de Münster : au sommet de la montagne Dülük Baba, voisine de la cité antique, ils ont découvert une stèle en basalte de 130 centimètres de hauteur représentant *Jupiter Dolichenus*. Le dieu lui-même est représenté debout sur un taureau, affublé de ses attributs habituels : une double hache et un faisceau d'éclairs ; face à lui, sa déesse partenaire est debout sur un cerf. En bas de la stèle, deux prêtres syriaques font un sacrifice.

Le style de la représentation est primitif. Sans doute montre-t-elle le dieu comme on le représentait depuis toujours dans la région, même si des éléments gréco-romains attestent que la stèle est d'époque romaine. Au sommet de Dülük Baba, les archéologues ont par ailleurs trouvé les preuves d'une utilisation religieuse du lieu dès le VI^e siècle avant notre ère. En particulier, ils ont retrouvé des centaines de milliers d'os d'animaux sacrifiés par le feu et, chose rare, plus d'un millier de sceaux et tampons aux motifs typiques du Proche-Orient de l'époque. De grande valeur, ces objets ont manifestement été sacrifiés par des personnes très pieuses. Parce qu'ils témoignent d'une très longue utilisation religieuse du sommet de Dülük Baba, ces indices suggèrent qu'un culte ancien et réputé y fut transformé à l'époque romaine en culte central de *Jupiter Dolichenus*.

F. S.

Le mystère des nuages nacrés

Le ciel des hautes latitudes est propice aux spectacles. Outre les aurores boréales, le promeneur peut parfois admirer au crépuscule des nuages nacrés, des nuages irisés qui apparaissent dans la partie supérieure de la stratosphère et dans la mésosphère. Le satellite AIM de la NASA a recueilli des informations sur leur formation. Analysées par James Russell, de l'Université Hampton, en Virginie, et ses collègues, ces données bousculent ce que l'on pensait savoir.

D'abord, les nuages nacrés ne sont pas limités à une altitude de 82 kilomètres comme on le pensait, mais à une bande plus large, de 79 à 90 kilomètres. Ensuite, ils sont dix fois plus brillants que prévu. Enfin, ils sont plus fréquents, plus étendus et l'on peut désormais les voir à des latitudes plus basses.

Ces modifications seraient dues à l'augmentation des quantités de gaz à effet de serre dans l'atmosphère, dont les principaux sont le méthane et le dioxyde de carbone. Le méthane s'oxyde en altitude et fournit de l'eau qui alimente les nuages nacrés. Et dans la haute atmosphère, le dioxyde de carbone diminue la température et favorise ainsi la naissance de ces nuages.

L. M.

AGU Meeting, San Francisco, 2007



V. Vénas

Terminus à puces

Votre chat a des puces et celles-ci vous envahissent ? Au lieu d'utiliser des produits toxiques qui ne font de bien à personne, passez plutôt un bon coup d'aspirateur, sinon à votre matou, du moins à vos tapis. Telle est la leçon des travaux de Fred Hink et Glen Needham, deux entomologistes américains de l'Université de l'État d'Ohio.

En 1996 et en 2000, d'autres scientifiques avaient montré que l'aspirateur retirait d'un tapis respectivement 90, 50, 63 et 95 pour cent des œufs, larves, pupes et adultes de puces de chat (*Ctenocephalides felis*) qui s'y trouvent. Ces indésirables sont retirés du tapis, certes, mais quel est leur sort ? Le sac de l'aspirateur est-il une simple escale avant qu'ils reprennent des forces, ou bien est-il pour eux un funeste terminus ?

La série de tests effectuée par F. Hink et G. Needham montre que l'aspirateur fait un massacre : 96 pour cent des puces adultes aspirées sont tuées, chiffre qui monte à 100 pour les larves et les pupes. Et des expériences de contrôle prouvent qu'une éventuelle toxicité des sacs d'aspirateur ou le simple écoulement d'air ne sont pas en cause : les puces et leurs stades non adultes, à la cuticule plus fragile, sont probablement écrasés contre les parois internes de l'aspirateur sous l'effet des pales, brosses et puissants flux d'air.

M. M.

Entomologia Experimentalis et Applicata, vol. 125 (2), pp. 221-222, 2007

ACTUELLEMENT EN KIOSQUE



Quel est le rôle de l'eau dans le système Terre ?
Où se trouvent nos ressources en eau douce ?
Celles-ci sont-elles bien gérées ?
Ce dossier fait le constat de leur fragilité
et envisage des solutions pour l'avenir.



Comment le cerveau reconnaît-il les visages sans
se tromper ? Comment les profileurs psychologiques
décodent-ils les personnalités politiques ? Pourquoi
jouons-nous au Loto ? Cerveau&Psycho répond
à ces questions et aborde d'autres sujets étonnants.

Au XIX^e siècle, Pasteur fait une découverte
redoutable : vers à soie, vin, bière, poules,
hommes... tous sont à la merci d'un même fléau,
les microbes. Revivez cette enquête passionnante
et les nombreuses controverses qui l'ont émaillée.



**POUR LA
SCIENCE**

Pour plus d'informations sur ces numéros,
rendez-vous sur notre site
www.pourlascience.com

BON DE COMMANDE

À retourner accompagné de votre règlement à **Pour La Science - Livres - BP 2 - 28410 Saint Lubin de la Haye**

Revue	Code	Prix unitaire	Quantité	Total
Les Dossiers	077658	6,90 €		
Les Génies de la Science	075092	6,90 €		
Cerveau & Psycho	076025	6,90 €		
Frais de port	France : 4,90 €	Étranger : 12 €		
Total commande				

Mes Coordonnées (à remplir) :

Nom, Prénom : _____

Adresse : _____

Code Postal : _____

Ville : _____ Pays : _____

E-mail* : _____

Téléphone* : _____ Date de naissance* : _____

Je règle par :

☐ Chèque à l'ordre de **Pour La Science**

☐ Carte bancaire

Numéro de la carte

Date d'expiration

Signature

En application de l'article 21 de la loi du 6 janvier 1978, les informations ci-dessus sont indispensables au traitement de votre commande. Elles peuvent donner lieu à l'exercice d'un droit d'accès et de rectification auprès de Pour La Science. Par votre intermédiaire, vous pouvez être amené à recevoir des propositions d'autres organismes. En cas de refus de votre part, il vous suffira de nous prévenir par simple courrier.



© Ron Miller

Mars : l'effet de serre sent le soufre

Le dioxyde de soufre serait responsable de la température clémente qui a régné sur Mars au début de son histoire, ainsi que de l'absence mystérieuse de carbonates à sa surface.

Il y a plus de 3,5 milliards d'années, la température était suffisamment clémente sur Mars pour que l'eau soit présente à l'état liquide à sa surface. Compte tenu de son éloignement du Soleil, un tel climat n'a pu exister que grâce à l'effet de serre. Quel était le responsable de cet effet ? Selon Itay Halevy, de l'Institut de technologie du Massachusetts, et ses collègues, il s'agirait du dioxyde de soufre.

On sait depuis peu que le dioxyde de carbone, en grande partie responsable de l'effet de serre sur Terre, n'en était pas la cause sur Mars. En effet, une concentration élevée de dioxyde de carbone dans l'atmosphère martienne primitive aurait dû entraîner la formation de carbonates. Mais aucune trace de ceux-ci n'a été décelée, alors même que la présence d'argiles implique que l'environnement était assez aqueux pour que ces roches se forment.

D'où viennent alors l'effet de serre et les abondants sulfates observés ? Le dioxyde de soufre serait la clef du mystère. Puissant gaz à effet de serre, il a sans doute été dégazé en abondance par le volcanisme au début de l'histoire martienne. Selon I. Halevy, il aurait régi la géochimie et le climat primitifs de Mars. En concentration élevée dans l'atmosphère, le dioxyde de soufre aurait ainsi maintenu une température clémente sur la planète rouge. Dissous dans l'eau, il aurait accaparé le magnésium et le calcium pour former des sulfites, empêchant ces éléments de participer à la formation de carbonates, sans toutefois rendre l'eau trop acide, ce qui autorise la formation d'argiles. En s'oxydant, les sulfites auraient ensuite donné naissance aux sulfates observés aujourd'hui. Puis, lorsque le volcanisme s'est tari, faute de gaz à effet de serre, Mars est devenue aride et froide.

Ce scénario diffère de celui en vigueur jusqu'ici, selon lequel le dioxyde de carbone était abondant, mais où les carbonates ont été détruits par l'acidification de l'eau due au soufre. Les futures données spectrométriques des sondes *Mars Express* et *Mars Reconnaissance Orbiter* permettront sans doute de trancher.

Ph. R.-G.

Science, vol. 318, pp. 1903-1907, 2007

Lucioles moins efficaces

Les lucioles (*Photinus pyralis*) produisent de la lumière par des réactions chimiques avec une efficacité extrêmement élevée. C'est du moins ce qu'on pensait depuis que ce rendement a été mesuré *in vitro*, en 1959. Mais Yoriko Ando, de l'Université de Tokyo, et ses collègues ont fait de nouvelles mesures avec des instruments beaucoup plus performants, et ont trouvé un rendement plus de deux fois inférieur.

La réaction luminescente, reproduite en laboratoire, met en jeu une molécule, la luciférine, qui s'oxyde en présence de luciférase (une enzyme), d'ATP (adénosine triphosphate) et d'ions de magnésium (Mg^{2+}). La luciférine oxydée est produite dans un état excité, et elle peut se désexciter en émettant un photon. Le « rendement quantique » de la réaction est la probabilité d'émission d'un photon pour une molécule de luciférine ayant réagi. L'ancienne mesure donnait un rendement maximal de 88 (± 25) pour cent, valeur qui s'est révélée par la suite douteuse. La nouvelle valeur est de 41,0 ($\pm 7,4$) pour cent, rendement qui reste élevé. Les chercheurs japonais ont également précisé l'influence de l'acidité du milieu sur les composantes spectrales (les couleurs) de l'émission lumineuse. M. M.

Nature Photonics, vol. 2, pp. 44-47, 2008

Le cerveau des plongeurs

Le 12 juillet 2007, Stéphane Mifsud a établi un record du monde d'apnée en restant immergé dix minutes et quatre secondes. Un autre mammifère, le phoque de Weddell, reste sans souci plus de 90 minutes sous l'eau. Comment expliquer un tel écart de performances ? Terrie Williams et ses collègues de l'Université de Californie, à Santa Cruz, ont montré qu'il tenait aux quantités de protéines (des globines) qui véhiculent l'oxygène dans le cerveau.

Chez 16 espèces de mammifères, les biologistes ont étudié non seulement l'hémoglobine (la globine circulante), mais aussi deux autres globines, découvertes en 2000, qui résident dans le cerveau, la cytoglobine et la neuroglobine. La première facilite le déplacement de l'oxygène vers les tissus nerveux ; la seconde empêche la formation de radicaux libres délétères. Les animaux plongeurs ont jusqu'à dix fois plus d'hémoglobine dans le cerveau que les animaux terrestres, et ont également plus de globines résidentes. Ces résultats contredisent l'hypothèse selon laquelle les prouesses des mammifères marins sont dues à une adaptation de leur système vasculaire. C'est peut-être aussi une piste à explorer pour protéger le cerveau en cas d'accident vasculaire. L. M.

Proc. R. Soc. B, prépublication, 2007



Chantal Rousselin / Palais de la Découverte

Les limites de la physiologie humaine

À haute altitude, tout un chacun a des maux de tête et des nausées, voire des œdèmes : c'est le mal des montagnes. L'intensité des symptômes varie selon les individus, mais comme l'explique Jean-Paul Richalet, de l'Hôpital Avicenne, à Bobigny, l'organisme humain a de grandes difficultés à s'adapter au manque d'oxygène.

Pour la Science : L'organisme humain a des capacités étonnantes d'adaptation et de récupération. Toutefois, il ne s'adapte pas sous l'eau et supporte mal la haute montagne. Sous l'eau, l'absence d'air est une limite dont il ne peut s'affranchir. Quelle est la principale limite en haute montagne ?

Jean-Paul Richalet : En fait, le corps humain a de nombreuses limites qui sont liées au froid, au chaud ou encore au stress. L'oxygène est une autre de ces limites et c'est elle qui est la cause principale des difficultés rencontrées en montagne : toutes les cellules du corps humain, consommant de l'oxygène, et les en priver entraîne de graves conséquences. La pression diminue avec l'altitude, mais ce n'est pas cela qui pose problème, ou plutôt indirectement : quand la pression atmosphérique baisse, la pression partielle d'oxygène diminue, c'est-à-dire que la quantité d'oxygène disponible pour l'organisme est réduite.

Pour la Science : À quoi sert l'oxygène dans le fonctionnement habituel ?

Jean-Paul Richalet : Comme pour beaucoup d'organismes vivants qu'on appelle aérobies – c'est-à-dire dont le métabolisme dépend de l'oxygène –, l'oxygène est nécessaire à la production de l'énergie indispensable au fonctionnement des cellules. Dès lors, quand l'approvisionnement en oxygène est réduit, l'organisme réagit pour compenser ce manque d'oxygène dans l'air : il essaye de faire en sorte qu'il en arrive toujours autant au niveau des mitochondries, ces petites usines cellulaires qui produisent de l'énergie à partir de l'oxygène. Deux organes sont sollicités : le cœur et les poumons. Le rythme cardiaque et la respiration accélèrent pour que le sang soit renouvelé plus rapidement dans les cellules. Ce sont deux réflexes immédiats du corps humain quand il est exposé à un manque

d'oxygène. De petites cellules localisées dans le cou, les chémorécepteurs, détectent le manque d'oxygène dans le sang, en informent le cerveau et, en quelques fractions de seconde, le cœur et la respiration s'accroissent.

Pour la Science : Est-ce que cette accélération du cœur et de la respiration parvient à compenser le déficit ?

Jean-Paul Richalet : En partie, mais on ne parviendra jamais par ces deux moyens à compenser le manque d'oxygène dans les alvéoles pulmonaires. Dès lors, on marche moins vite, on est essoufflé, on doit s'arrêter souvent, on est plus vite fatigué. Dans un premier temps, l'organisme subit ce manque d'oxygène : le rythme de la respiration et celui du cœur restent élevés. Mais dans un deuxième temps, au bout d'une semaine environ, il essaye de compenser le manque d'oxygénation en fabriquant davantage de globules rouges : la moelle osseuse commence à produire plus de globules rouges, qui transportent plus d'oxygène aux cellules, notamment aux muscles et au cerveau. Cela correspond à un début d'adaptation.

Pour la Science : En une semaine, si l'on est passé de 1 000 à 3 000 mètres, tout redevient-il normal ?

Jean-Paul Richalet : Malheureusement non. Mais on se sent mieux, c'est-à-dire que pendant cette semaine, on a passé le stade du mal des montagnes : on est moins essoufflé, mais on ne récupère pas la performance physique qu'on avait au niveau de la mer. Et ce quelle que soit la durée : que l'on reste deux ans, dix ans ou 20 ans, on est toujours diminué physiquement en altitude. Bien que l'on ait plus de globules rouges et que la respiration soit plus efficace, on est moins performant. Bien sûr, on peut survivre ce qui montre que l'organisme humain dispose de mécanismes d'adaptation au manque d'oxygène.

Pour la Science : Qu'est-ce que le mal des montagnes ?

Jean-Paul Richalet : C'est une traduction du fait que les mécanismes d'adaptation du corps au manque d'oxygène ne sont pas encore efficaces. Donc, c'est une sorte de signal d'alarme indiquant que l'organisme n'est pas encore adapté au manque d'oxygène. On n'a pas le mal des montagnes quand on monte à l'Aiguille du Midi, et que l'on y reste une heure avant de redescendre. Le mal des montagnes apparaît au bout de quatre à six heures. Ainsi, quand les touristes arrivent à Cuzco, au Pérou, à 3 400 mètres d'altitude, ils se sentent bien. Mais quand vient le soir, ils commencent à avoir mal à la tête et des nausées. Ils dorment mal, n'ont plus faim, se sentent fatigués, sont essoufflés, vomissent parfois. Ils ressentent une sensation de malaise général.

Pour la Science : Et là encore, le manque d'oxygène est en cause ?

Jean-Paul Richalet : Le manque d'oxygène perturbe l'agencement des membranes des vaisseaux sanguins, qui deviennent perméables. Ils se mettent à fuir, ce qui crée des œdèmes. Parfois, c'est le visage qui enfle, parfois du liquide s'accumule dans les poumons et l'on a des difficultés à respirer. L'œdème peut se localiser dans le cerveau : le liquide qui s'accumule dans la boîte crânienne comprime les structures cérébrales, ce qui entraîne rapidement une

perte de connaissance, voire la mort si on ne redescend pas rapidement. Si la descente est suffisamment rapide, l'œdème pulmonaire régresse et l'on récupère sans séquelles. En ce qui concerne un œdème cérébral, on constate souvent de petits troubles de la mémoire pendant quelques semaines, voire quelques mois, parfois des troubles dépressifs, mais la récupération finit par être totale. Qui plus est, au-delà de 3 000 à 3 500 mètres, quasiment tous les individus sont touchés. L'intensité des symptômes peut varier, mais c'est exceptionnel que l'on n'ait pas un peu mal à la tête ou quelques nausées. Il y a même des gens qui souffrent du mal des montagnes en allant faire du ski dans des stations moins élevées ! Toutefois, on peut réduire l'intensité des symptômes en laissant à l'organisme le temps de s'acclimater. Ainsi, il est conseillé, à partir de 3 000 mètres, de ne pas dépasser 400 mètres de dénivelé entre deux nuits consécutives. Dans la journée, on peut passer un col et dépasser le dénivelé, ce qui compte c'est la différence d'altitude entre deux nuits. Au bout d'un mois environ, l'organisme est acclimaté.

Pour la Science : Que se passe-t-il à très haute altitude, par exemple à 5 000 mètres ?

Jean-Paul Richalet : Jusqu'à 5 000 mètres, on peut encore s'adapter, c'est-à-dire qu'on va produire plus de globules rouges, on va respirer un peu plus vite. En revanche, le



© Shutterstock/Pavel Kmeto

cœur va ralentir pour essayer d'économiser un peu d'énergie. Tous ces mécanismes font que l'on réagit encore correctement au manque d'oxygène. Au-delà, on entre dans un domaine qu'on pourrait qualifier d'extraphysiologique : on ne peut pas vivre de façon permanente au-delà de 5 500 mètres. On perd progressivement des muscles et du poids, on se déshydrate, on perd aussi certainement des neurones, car ces cellules sont très sensibles au manque d'oxygène. Le sommeil est plus léger. On a aussi des apnées du sommeil, c'est-à-dire des arrêts de la respiration pendant le sommeil. L'appétit diminue et les nausées sont fréquentes. Les centres nerveux qui contrôlent la faim sont perturbés, donc on mange moins et on perd du poids. Lors des premières grandes expéditions dans l'Himalaya, les

Les populations qui vivent en altitude ont développé divers mécanismes d'adaptation : une meilleure respiration, une plus grande efficacité des poumons ou une surproduction de globules rouges.

gens perdaient parfois 10 à 15 kilos. Et si l'on ne redescend pas rapidement (au bout de un, maximum deux mois), les dégradations peuvent être irréversibles. La plupart des camps de base sont situés aux alentours de 5 000 mètres. Le camp de base de l'Everest est à 5 400. Les accidents sont fréquents, notamment des œdèmes pulmonaires ou cérébraux. Les alpinistes montent rapidement en partant de 3 000 mètres, sans s'acclimater. Ils arrivent au camp de base et, au début, ils se sentent bien, voire euphoriques. Pourtant, un œdème pulmonaire peut survenir dès la première nuit au camp de base. Beaucoup récupèrent à peu près. Plus un tel accident survient à haute altitude, plus il est difficile de récupérer. Il devient quasiment impossible de dormir à 8 000 mètres. À cette altitude, on peut survivre, mais on ne peut pas vivre. Il ne faut pas y rester trop longtemps. Typiquement, pour l'ascension de l'Everest, il faut à peu près un mois. Une grande partie se passe au camp de base et il faut limiter le nombre de nuits passées au-delà de 6 000 mètres à une dizaine.

Pour la Science : Est-ce que l'utilisation de bouteilles d'oxygène permet d'améliorer les performances ?

Jean-Paul Richalet : Cela les améliore, mais on ne récupère pas des capacités normales, car on respire de l'air seulement un peu enrichi en oxygène. La première ascension de l'Everest sans oxygène a été réalisée en 1978 par l'Italien Reinold Messner et l'Autrichien Peter Habeler. Depuis, il y a à peine plus d'une centaine de personnes qui ont réussi à faire l'ascension sans oxygène, et beaucoup sont morts en faisant cette tentative. Donc, on est vraiment à la limite de la physiologie humaine, de ce que peut tolérer l'homme en termes de manque d'oxygène.

Pour la Science : Pourtant, certaines populations des Andes, par exemple, vivent à plus de 3 000 mètres. Ont-elles des mécanismes d'adaptation particuliers ?

Jean-Paul Richalet : On pourrait le croire en regardant par exemple les sherpas qui portent des charges de 30 à 40 kilogrammes sur le dos et qui grimpent jusqu'à 8 000 mètres sans oxygène. Pourtant, on n'a pas trouvé le « gène de l'altitude ». Il n'y a pas d'adaptation génétique à l'altitude. Ils ont la même hémoglobine, les mêmes globules rouges, les mêmes mécanismes de contrôle de la respiration. Les personnes qui vivent depuis leur naissance en altitude développent plutôt des mécanismes d'acclimatation. On a mis en évidence plusieurs stratégies. Contrairement aux Andins qui ont beaucoup de globules rouges pour améliorer le transport de l'oxygène, les sherpas, au Tibet et au Népal n'en ont pas beaucoup. En revanche, ils respirent mieux. Quant aux Éthiopiens et aux Kenyans, leurs mécanismes de diffusion de l'oxygène au niveau des poumons sont plus performants. Notons que l'adaptation des Andins (l'augmentation du nombre de globules rouges) peut devenir un inconvénient : entre 10 à 15 pour cent des personnes qui vivent sur les hauts plateaux péruviens et boliviens finissent par avoir trop de globules rouges quand ils vieillissent. Ils ne respirent pas assez, de sorte que leur sang s'appauvrit en oxygène, ce qui stimule la production de globules rouges. Le sang devient extrêmement visqueux, entraînant des thromboses et des accidents vasculaires cérébraux. C'est une anomalie du contrôle de la respiration. Ce problème de santé publique touche des millions d'individus.

Pour la Science : Est-ce pour produire plus de globules rouges que les sportifs s'entraînent parfois en altitude ?

Jean-Paul Richalet : Quand un sportif va s'entraîner à la montagne, ce n'est pas pour s'oxygéner, puisqu'il y a moins d'oxygène en montagne qu'à basse altitude. En revanche, un séjour en altitude stimule la production des globules rouges. Quand il revient au niveau de la mer, il a plus de globules rouges qui transportent plus d'oxygène ; il espère ainsi améliorer ses performances. En fait, la production n'augmente que si l'on reste au moins deux ou trois semaines, à plus de 2 500 mètres. Malheureusement comme la performance physique diminue en altitude, les sportifs ne peuvent pas s'entraîner au niveau qui est le leur à basse altitude, car ils sont essouffés. Ainsi, la production de globules rouges augmente, mais l'entraînement est de moindre qualité. Aujourd'hui, certains athlètes vont dormir en altitude pour fabriquer des globules rouges pendant leur sommeil, et redescendent dans la journée pour s'entraîner à basse altitude pour que leur performance physique soit normale. Certains utilisent une autre technique : ils installent chez eux une tente sous laquelle ils font circuler un mélange d'air appauvri en oxygène : ils dorment chez eux en altitude simulée, mais, au matin, ils respirent l'air ambiant. Est-ce une forme de dopage ? Pour le moment, le Comité olympique international ne l'interdit pas, parce que cela reproduit les conditions qui règnent en altitude, mais cette pratique soulève des questions éthiques. Quelle que soit la méthode utilisée, l'organisme constate que la quan-

tité de globules rouges est trop élevée par rapport à la concentration d'oxygène, de sorte que ces globules rouges disparaissent plus vite que la normale, en trois à quatre semaines, alors que la durée de vie moyenne d'un globule rouge est plutôt de quatre mois. Tous les globules rouges en excès disparaissent rapidement.

Pour la Science : Les alpinistes sont confrontés au manque d'oxygène, mais aussi au froid. S'adapte-t-on au froid ?

Jean-Paul Richalet : Non, le froid est une autre limite du corps humain. Nous sommes des animaux à sang chaud, c'est-à-dire que nous nous acclimatons plus facilement au chaud qu'au froid. Les animaux bien adaptés au froid ont une fourrure, dont les poils peuvent se dresser pour emmagasiner de l'air, ce qui fait une couche isolante. Nous, notre seule façon de résister au froid est de frissonner. Le frisson est une contraction involontaire des muscles localisés sous la peau, et ces contractions produisent de l'énergie. Dès que les muscles se contractent, 80 pour cent de l'énergie sont dissipés sous forme de chaleur et 20 pour cent seulement sont transformés en énergie mécanique. Pourtant, on ne résiste pas longtemps au froid grâce à l'énergie dégagée par les frissons : cette production de chaleur est peu efficace. Les alpinistes confrontés au froid peuvent être soumis à deux contraintes : les gelures ou l'hypothermie (c'est la température du corps entier qui baisse, par exemple de 37 °C à 36,5). Pour lutter contre l'hypothermie, les vaisseaux périphériques se contractent. La circulation se réduit de façon à continuer à alimenter les organes vitaux, mais c'est au détriment des extrémités. En fait, pour résister au froid, l'homme ne peut mettre en jeu que des stratégies comportementales ou des moyens technologiques : vêtements, boissons chaudes, chauffage.

Pour la Science : Peut-on éviter le mal des montagnes ?

Jean-Paul Richalet : Pour réduire les risques de se sentir mal et favoriser une bonne acclimatation, il faut prendre son temps. Ce conseil s'applique à tous, mais nous ne sommes pas égaux devant l'hypoxie : certains s'adaptent mieux que d'autres. Avant de partir, on peut faire un test d'effort en hypoxie pour savoir si l'on est ou non sensible au mal des montagnes. Durant ce test, la personne doit pédaler sur une bicyclette et on lui fait respirer un mélange d'air appauvri en oxygène. En mesurant la respiration, le rythme cardiaque, le taux d'oxygène dans le sang, le médecin peut dépister les individus particulièrement sensibles. Les cellules qui, normalement, détectent le manque d'oxygène sont peu performantes, de sorte que ces personnes ne respirent pas assez vite et risquent de mal s'adapter à l'altitude. Il existe un médicament, l'acétazolamide, qui stimule le réflexe de la respiration. Il est assez efficace pour prévenir le mal des montagnes et favoriser l'acclimatation. Qui plus est, c'est une substance un peu diurétique, ce qui permet de réduire un éventuel œdème.

Pour la Science : L'acétazolamide pourrait-elle être utilisée chez les populations qui vivent en altitude et produisent trop de globules rouges, car leurs capteurs d'oxygène sont déficients ?

Jean-Paul Richalet : Oui, et notre groupe de recherche a été le premier à montrer que ce médicament peut effectivement être utilisé chez ces populations qui présentent une désadaptation à l'altitude. On a démontré qu'il améliore la respiration de ces personnes, de sorte que le nombre de leurs globules rouges diminue. La respiration étant meilleure, leur sang est mieux oxygéné et, par conséquent, la production de globules rouges diminue. Nous avons récemment achevé une étude qui a duré six mois dans la ville de Cerro de Pasco, au Pérou. C'est une ville minière située à 4 350 mètres d'altitude, où vivent 80 000 personnes. Or 10 à 15 pour cent de la population présentent un excès de globules rouges. L'administration d'acétazolamide normalise la quantité de globules rouges. De plus, ce médicament ne coûte pas cher. Avec nos collègues péruviens, nous essayons de développer un programme de santé publique où les populations qui vivent à très haute altitude pourraient bénéficier de ce traitement.

À haute altitude, tous les individus sont sujets au mal des montagnes. Seule l'intensité des symptômes change. Pour réduire au maximum les désagréments, la meilleure prescription est de prendre son temps pour faire l'ascension.

Pour la Science : Est-ce que la bonne condition physique favorise l'adaptation à l'altitude ?

Jean-Paul Richalet : Il faut être en forme physiquement, mais ce n'est pas une garantie contre le mal des montagnes. Même les athlètes les mieux entraînés y sont sensibles. Qui plus est, plus un athlète est entraîné, plus ses performances physiques diminuent avec l'altitude. Il n'y a pas de préparation physique particulière, la principale étant, comme nous l'avons rappelé, de prendre son temps. Par exemple, le Kilimandjaro est une destination assez prisée des personnes qui apprécient l'altitude. C'est un sommet à 5 900 mètres qui est très facile. En trois ou quatre jours, on franchit 4 000 mètres de dénivelé : tout le monde est malade ! Sauf ceux qui marchent lentement et qui s'arrêtent régulièrement. Le proverbe italien *Chi va piano va sano e chi va sano va lontano* (Qui va lentement va sûrement et va qui va sûrement va loin) est la règle essentielle.

Jean-Paul RICHalet, professeur de physiologie à l'Université Paris 13, dirige le Service physiologie, explorations fonctionnelles, médecine du sport du groupement hospitalier universitaire Nord-Avicenne.

Ce texte est un résumé d'une conférence que Jean-Paul Richalet a donnée au Palais de la Découverte. L'enregistrement de la conférence sera en ligne sur le site de *France Culture* (http://www.radiofrance.fr/chaines/france-culture/nouveau_prog/connaissance/alacarte.php) à partir du mois de février 2008.

Comment stocker le sang de cordon ombilical ?

Le potentiel des cellules souches contenues dans le sang de cordon ombilical est immense. Ce sang est stocké soit dans des banques publiques de qualité contrôlée et faisant partie d'un réseau international, soit dans des banques privées à but lucratif. Comment concilier ces deux types de stockage ?

Le sang contenu dans le placenta était jusqu'à une période récente considéré comme un déchet opératoire, et l'idée de l'utiliser pour traiter des maladies graves du sang est apparue il y a seulement 20 ans. En effet, le sang du nouveau-né contient des cellules étonnantes de vigueur, de capacité de régénération et d'immaturité immunologique. Ces propriétés en font une ressource unique de cellules souches réparatrices.

La première greffe de sang de cordon ombilical a été faite il y a près de 20 ans chez un enfant atteint d'une maladie héréditaire rare et mortelle de la moelle osseuse, la maladie de Fanconi. Son seul espoir de guérison était une greffe de moelle osseuse, traitement déjà appliqué avec succès dans notre Service à l'Hôpital Saint-Louis à Paris. Une greffe de moelle consiste à restaurer le stock des cellules sanguines d'un malade grâce à des cellules prélevées sur un donneur (le stock est détruit par diverses maladies évoquées ultérieurement). La moelle est prélevée chez le donneur sous anesthésie générale au niveau des os du bassin et injectée par voie intraveineuse au receveur. Après la transplantation, les cellules saines du donneur finissent par remplacer celles du receveur.

Les premières greffes se faisaient à partir d'un donneur – le plus souvent un frère ou une sœur – à condition qu'il (elle) ait le même système de compatibilité tissulaire ou système HLA (*Human Leukocyte Antigen*). En effet, les cellules de l'organisme portent des antigènes du soi – les antigènes HLA – qui reconnaissent comme étranger tout tissu issu d'une autre personne. Dans le cas des greffes de moelle, un risque supplémentaire existe : celui d'une réaction du greffon contre l'hôte. La moelle contient des cellules à divers stades de différenciation : depuis les cellules souches hématopoïétiques qui ne sont pas encore différenciées et donneront tous les types de cellules du sang, jusqu'à des cellules qui sont déjà engagées dans une voie de différenciation, c'est-à-dire qui ont déjà acquis une spécificité fonctionnelle. Ainsi, lors d'une greffe de cellules de moelle, on injecte

des cellules responsables de la protection de l'organisme contre les éléments étrangers, notamment des lymphocytes qui attaquent les cellules du receveur si les systèmes HLA sont trop différents. Pour éviter les rejets, on pratique une irradiation ou l'on administre une chimiothérapie ou des immunosuppresseurs, telle la ciclosporine, qui inhibent les réactions de défense de l'organisme, mais, surtout, on choisit un donneur dont le système HLA est le plus proche possible de celui du receveur.

Pourtant, même entre frères et sœurs, les plus proches sur le plan immunologique, une telle compatibilité tissulaire n'existe que dans un cas sur quatre. L'idée d'utiliser le sang du cordon ombilical à la place de la moelle osseuse est venue des travaux réalisés par l'équipe de Hal Broxmeyer, de l'Institut Rockefeller à New York, qui recherchait une source de cellules souches hématopoïétiques facilement stockables que l'on pourrait utiliser en cas d'accident nucléaire. C'est à cette époque que l'accident de Tchernobyl venait de survenir et que le risque nucléaire était dans tous les esprits. Ce chercheur a donc commencé par regarder si, dans le sang du cordon ombilical prélevé à la naissance, on retrouvait les cellules souches du sang : elles y étaient présentes en grand nombre, bien plus que dans le sang de donneurs adultes.

La preuve de leur efficacité clinique devait alors être apportée par Arleen Auerbach, de l'Université Rockefeller, à New York, qui travaillait sur la maladie de Fanconi et disposait d'échantillons de sang du fœtus pour le diagnostic prénatal de cette maladie. Elle a ainsi pu identifier plusieurs familles pour lesquelles on savait avant la naissance que l'enfant était indemne de la maladie génétique et qu'il avait un système HLA compatible avec celui du malade. Le sang du cordon ombilical a donc été collecté à la naissance et conservé dans l'azote à -180°C , puis utilisé après décongélation pour greffer l'enfant malade. Nous avons pratiqué cette greffe à l'Hôpital Saint-Louis, et elle a parfaitement réussi. Le patient a aujourd'hui plus de 20 ans. Il est guéri et mène une vie normale, ce qui démontre que le sang d'un seul cordon, repré-

sentant à peine 100 millilitres, peut reconstituer à long terme tout le système hématopoïétique et immunologique d'un individu. Cette première greffe a suscité beaucoup d'intérêt et d'interrogations sur les usages potentiels de cette nouvelle source de cellules souches. Comment profiter au mieux de cette ressource ? Comment la stocker ? Comment mettre les ressources mondiales en commun afin d'augmenter les chances de trouver des cellules dont les systèmes HLA sont les plus compatibles possible ? Et enfin, s'est posée – et se pose encore – la question de la nature des banques stockant le sang de cordon : les banques privées ont-elles leur place à côté des banques publiques ?

La greffe de sang de cordon : plus efficace que celle de moelle

Depuis la première greffe de sang de cordon ombilical, plus de 10 000 patients ont reçu ce type de greffe et des banques de sang de cordon ombilical se sont développées dans le monde entier : plus de 250 000 unités sont stockées. Les avantages de ce type de greffe sont nombreux : le prélèvement est sans danger pour la mère et pour l'enfant. Le placenta étant considéré comme un déchet opératoire, il est habituellement jeté, et le conserver ne présente aucun problème d'ordre éthique. Le prélèvement peut être conservé indéfiniment dans l'azote liquide. De plus, le risque de transmission d'infection est nul, car les greffons sont testés avant leur stockage. De surcroît, chez le nouveau-né, le greffon ne doit pas avoir obligatoirement le même système HLA, car le système immunitaire est immature et supporte des différences. Enfin, le greffon est immédiatement disponible, ce qui permet de ne pas attendre lorsque la greffe doit être faite en urgence, comme dans le cas d'une leucémie aiguë. Tous ces facteurs expliquent le développement rapide des banques de sang de cordon ombilical et l'augmentation du nombre de greffes ; dans certains pays, ce nombre est aujourd'hui supérieur à celui des greffes de moelle osseuse.

Les maladies traitées sont des maladies malignes, telles que les leucémies aiguës et chroniques, les lymphomes, les myélomes, et des maladies non malignes ou héréditaires, notamment la thalassémie, la drépanocytose ou anémie à cellules falciformes, certaines aplasies médullaires (des dysfonctionnements de la moelle), des déficits immunitaires et des maladies métaboliques héréditaires. L'analyse des résultats de ces greffes a révélé l'un des éléments qui ont contribué à l'augmentation des greffes de sang de cordon : une greffe faite à partir d'un sang de cordon non HLA identique donne des résultats aussi bons, si ce n'est meilleurs qu'une greffe de moelle prélevée chez un adulte ayant le même système HLA.

Par ailleurs, les divers résultats ont montré que si le nombre de cellules dans un cordon est insuffisant, on peut greffer

simultanément le sang de deux cordons même s'ils n'ont pas le même système HLA. Ainsi, le sang de cordon peut être considéré comme compatible avec tous les receveurs : c'est un sang de « donneur universel ». Pourquoi les différences des systèmes HLA n'empêchent-elles pas le succès de la greffe ? Parce que le sang administré contient moins de lymphocytes responsables des réactions de rejet et que ces cellules sont moins avancées dans la voie de la différenciation ; les lymphocytes sont donc moins « réactifs » que ceux des adultes.

Avant l'utilisation du sang de cordon, la moitié seulement des patients en attente de greffe de moelle pouvaient être greffés faute de donneurs compatibles, même si 11 millions de personnes sont inscrites sur les registres internationaux des donneurs de moelle. Une étude récente faite aux États-Unis a montré que sur 44 000 personnes à la recherche d'un donneur de moelle, 40 pour cent n'ont pas été greffés soit parce qu'ils n'avaient pas de donneur identique, soit parce que le temps d'attente pour trouver un donneur avait été trop long et qu'ils avaient rechuté ou qu'ils étaient décédés avant la greffe.



1. Le sang de cordon est prélevé au moyen d'une aiguille (en haut à droite). On récupère ainsi environ 100 millilitres de sang riche en cellules souches. Après avoir subi divers contrôles de qualité, le sang est stocké dans des poches (en haut à gauche) et congelé.

En ce qui concerne les greffes de moelle osseuse, 25 pour cent seulement des malades bénéficient d'un donneur ayant le même système HLA dans leur famille. La greffe de sang de cordon s'adresse donc aux patients pour lesquels il n'existe pas de donneur de moelle compatible ou parce que le donneur n'est pas disponible dans les délais nécessaires. On estime à 40 pour cent la chance de trouver un donneur de moelle osseuse dont le système HLA est compatible avec celui du malade : par conséquent, augmenter indéfiniment le nombre de donneurs de moelle est peu efficace. Au contraire, sur un total de 275 000 sangs de cordon présents dans les banques internationales, plus de 10 000 greffes ont déjà pu être réalisées, soit un taux d'utilisation de la ressource compris entre deux et trois pour cent. Il convient d'ajouter qu'une étude prenant en compte le coût des deux techniques accentuerait encore l'écart en matière d'efficacité au profit des greffes de sang de cordon. Qui plus est, l'efficacité de la transplantation est en faveur du sang de cordon : il suffit d'injecter 100 millilitres de sang de cordon, alors qu'il faut un litre de sang issu de la moelle, car le sang de cordon est beaucoup plus riche en cellules souches.

Deux types de banques

Il existe actuellement deux types de banques de sang de cordon ombilical : les banques dites allogéniques (le donneur est différent du patient ; il peut être apparenté ou totalement étranger) et les banques dites autologues qui réalisent un stockage pour le propre usage de l'enfant : c'est une sorte d'assurance biologique en cas de maladie ultérieure.

Dans le premier cas, le don est gratuit, altruiste et anonyme, et destiné à traiter un patient qui aurait besoin d'une telle greffe. Dans le second, il s'agit d'un usage privé sans indication thérapeutique avérée (on ignore si l'enfant en aura besoin et si le sang stocké pourra servir à traiter une éventuelle maladie). Les banques de sang de cordon allogéniques, le plus souvent publiques, collectent le sang de cordon dans des maternités agréées où une information préalable est faite aux femmes enceintes afin qu'elles donnent un consentement éclairé. Toutes les précautions sont prises pour assurer la bonne qualité du greffon et seulement la moitié des cordons prélevés remplissent les critères pour être intégrés dans la banque. Une organisation internationale met à disposition par informatique la liste de tous les cordons disponibles dans les banques, de telle façon que les médecins greffeurs puissent à tout moment faire la recherche du meilleur greffon pour un patient donné.

L'organisation *Netcord Eurocord* a été fondée en 1995. *Netcord* est un réseau de banques de sang de cordon, qui établit les standards, réalise des contrôles de la qualité des prélèvements et recherche de nouveaux donneurs. Quant à *Eurocord*, cette institution est responsable du volet clinique (enregistrement et analyse des cas, rédaction de protocoles, réalisation des études cliniques). *Eurocord* travaille en étroite collaboration avec les banques, les centres de greffes et les registres de donneurs de moelle. *Netcord Eurocord* a pour mission de collecter les unités de sang de cordon ombilical, de les conserver au froid et de fournir des unités ayant les meilleurs critères de qualité et de sécurité. Il doit aussi amé-

liorer la probabilité de trouver un donneur de cellules souches hématopoïétiques pour tous les patients, quelle que soit leur origine ethnique et géographique, promouvoir l'utilisation clinique du sang de cordon ombilical, notamment tenir un registre international des ressources. Enfin, il doit aussi promouvoir la recherche sur la biologie et les applications cliniques du sang de cordon ombilical.

En 2006, 147 000 unités de sang de cordon étaient stockées aux États-Unis, 92 000 en Europe, 32 000 au Japon, 17 000 en Australie, 1 480 en Amérique du Sud. Le réseau *Netcord Eurocord* dispose désormais de 154 000 unités de sang de cordon accréditées, dont 35 200 à New York, 12 500 à Düsseldorf, 12 400 à Milan et 6 000 en France. Le réseau *Netcord Eurocord* a contribué à près de 6 000 greffes. Le nombre de greffes a régulièrement augmenté depuis une dizaine d'années (voir la figure 2).

Qui plus est, ce sont aujourd'hui en majorité des adultes qui bénéficient de ces greffes : la méthode, d'abord utilisée chez le jeune enfant, concerne désormais plus d'adultes que d'enfants. En Europe, le nombre de greffes de sang de cordon a fortement augmenté après la publication en 2004 de deux études donnant les mêmes résultats en Europe et aux États-Unis, et montrant que chez l'adulte atteint de leucémie, la greffe de sang de cordon ombilical n'ayant pas le même système HLA donne les mêmes résultats que la greffe de moelle HLA compatible.

Un réseau français en retard

En France, le dispositif est géré par l'Agence française de biomédecine dans le cadre d'un réseau comprenant l'Établissement français des greffes et *France greffe de moelle*. Deux banques sont implantées à Besançon et à Bordeaux. D'autres le seront bientôt à Paris et à Marseille. Une trentaine de centres pratiquent les transplantations. L'objectif consiste à produire et stocker au moins 10 000 unités de sang de cordon en cinq ans. Cependant, nous devrions disposer de neuf unités pour 100 000 habitants, soit d'un total de 50 000 unités, pour atteindre le niveau des banques internationales. Ainsi, la France a pris du retard dans le développement des banques de sang de cordon ombilical. Pourtant, elle est l'un des pays les plus actifs en matière de greffes, ce qui entraîne une importation d'environ 70 pour cent des cordons utilisés. Le nombre de greffes de sang de cordon devrait doubler grâce à l'augmentation du nombre de banques de sang de cordon ombilical, la possibilité d'utiliser des donneurs non HLA identiques et de pratiquer, si nécessaire, une double greffe. Cette augmentation envisagée du nombre de greffes de sang de cordon constitue un réel enjeu de santé publique.

Quelle est la situation dans les autres pays ? Aux États-Unis, le nombre de greffes de sang de cordon constitue une activité soutenue qui progresse régulièrement au détriment des greffes de moelle osseuse. Par ailleurs, le *National Cord Blood Program* (programme national de sang de cordon) a été mis en œuvre. Le Congrès a lancé une étude afin de déterminer comment développer l'utilisation du sang de cordon ombilical. Son rapport, rédigé en 2005, a confirmé le fait que le sang de cordon constitue une source précieuse de cellules souches. Il a été demandé

un soutien fédéral pour les banques publiques de sang de cordon, ainsi qu'une coordination avec le programme relatif aux donneurs de moelle. Une structure de coordination a été mise en œuvre, et le Congrès américain a attribué 25 millions de dollars par an durant quatre ans.

Le Japon s'est également investi dès le début dans les programmes de banques de sang de cordon. Ce pays dispose désormais de 11 banques et de 220 centres de transplantation. Le nombre de greffes de sang de cordon y est désormais plus important que le nombre de greffes de moelle.

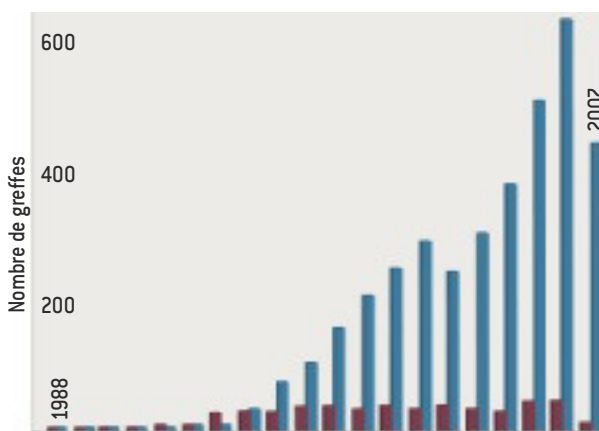
Quelle place pour les banques privées ?

Nous l'avons évoqué, les banques publiques (allogéniques) présentent de nombreux avantages en termes de garantie des ressources, de disponibilité, de concertation grâce au fichier international. Dans ces conditions, quelle est la place de la conservation privée (autologue) ? Un certain nombre d'arguments peuvent être avancés pour justifier la conservation autologue de sang de cordon ombilical : il n'existe pas d'interdiction de conserver ses propres cellules ; il s'agit d'une chance unique de conservation au titre d'une « assurance biologique », garantie contre les risques éventuels de maladie à venir ; les cellules peuvent être conservées indéfiniment (les cellules congelées à -180°C il y a 20 ans ont gardé leurs propriétés).

Il convient de rappeler qu'après dix ans de stockage dans une banque publique, un patient en attente de greffe de moelle a 35 pour cent de chances de trouver un donneur familial, 70 de trouver une unité de sang de cordon ou un donneur de moelle dans un registre international, mais quasiment 100 pour cent de chances de récupérer son sang de cordon.

Plusieurs modèles hybrides permettant de concilier les deux options ont été proposés : la coexistence de deux dispositifs permet à la mère de choisir de stocker le sang de cordon dans une banque autologue ou dans une banque allogénique ; on peut conserver le sang sur un mode autologue, mais le céder au public en cas de besoin ; chaque unité de sang de cordon peut être partagée entre la banque autologue (20 pour cent), et la banque allogénique (80 pour cent), mais ce modèle présente l'inconvénient de réduire une quantité de cellules déjà parfois insuffisante pour une greffe.

La comparaison entre banques allogéniques et banques autologues fait apparaître que si les bénéfices sont avérés pour les banques allogéniques, ils ne le sont pas pour les banques autologues ; la qualité des banques allogéniques est validée, celle des banques autologues est variable ; les banques allogéniques sont des banques publiques ne dégagant pas de profit, contrairement aux banques autologues privées ; le suivi est garanti pour les premières, pas pour les secondes ; les banques allogéniques reposent sur le principe de solidarité, les banques autologues sur un modèle de compétition avec les banques allogéniques dans le partage des ressources ; l'égalité d'accès est assurée dans les banques allogéniques seulement ; l'information est ouverte dans les premières, mais peut être biaisée dans les secondes ; enfin, seules les banques allogéniques font l'objet d'un inventaire mondial.



2. Évolution du nombre de greffes (apparentées en rouge et non apparentées en bleu) de sang de cordon par an dans le registre Eurocord, depuis 1988, impliquant 510 centres de greffe et 45 banques dans le monde.

Pour sa part, Netcord a pris une position nuancée. Il soutient l'établissement de banques publiques, fondées sur des dons altruistes et volontaires. Il insiste sur le fait que les familles doivent recevoir une information impartiale concernant les bénéfices et les risques du stockage privé afin d'émettre un consentement éclairé. Il a également demandé le respect de critères de qualité identiques concernant les deux types d'établissements, alors que pour l'instant les banques privées exercent en l'absence totale de transparence. Netcord souhaite aussi que la promotion et le financement des banques autologues – en l'absence d'indication médicale – ne relèvent pas de l'argent public.

Outre l'intérêt que présentent les cellules souches de sang de cordon dans diverses indications thérapeutiques connues, il existe une facette du potentiel de ces cellules que nous n'avons pas évoquée : la médecine régénérative. Les cellules souches totipotentes ont un potentiel de différenciation quasi illimité, c'est-à-dire qu'il est possible de faire apparaître l'ensemble des catégories de cellules à partir du sang de cordon. On pourrait donc les utiliser pour remplacer des cellules malades de l'organisme. Le sang de cordon ombilical constitue une des sources possibles de cellules souches qui présentent plusieurs avantages par rapport aux autres sources : les cellules embryonnaires ou les cellules souches adultes. Les cellules souches de cordon présentent un potentiel de différenciation notablement supérieur à celui des cellules souches adultes ; leur nombre est supérieur ; leur isolement est plus simple ; enfin, elles ne posent pas les problèmes éthiques soulevés par les cellules souches embryonnaires.

Rappelons que les banques privées sont extrêmement rentables, car elles ne mènent aucune recherche et ne font que stocker les prélèvements. Or compte tenu du potentiel thérapeutique que représentent ces cellules, et des recherches à mener pour préciser ce potentiel, il convient de se demander si les avancées de la recherche justifient la conservation autologue du sang de cordon.

Éliane GLUCKMAN est responsable du programme Eurocord à l'Hôpital Saint-Louis, à Paris.

Ovnis : des enquêtes sujettes à caution

Les investigations du service français d'étude des phénomènes aérospatiaux non identifiés, le GEIPAN, manquent de rigueur. Cela sème un doute sur l'ensemble des travaux relatifs aux ovnis.

Le 12 novembre dernier, à Washington, 19 personnalités – anciens pilotes, scientifiques et responsables politiques – lançaient un appel pour réclamer la réouverture d'une enquête officielle sur les ovnis. Depuis le début de la médiatisation du phénomène aux États-Unis en 1947, des dizaines de milliers d'observations d'objets volants non identifiés ont été rapportées à travers le monde. De multiples commissions ont été constituées pour étudier et tenter d'expliquer ces cas, sans parvenir à des conclusions totalement convaincantes. En France, c'est la mission d'un service du Centre national d'études spatiales (CNES), le groupe d'étude et d'information sur les phénomènes spatiaux non identifiés, ou GEIPAN. Créé en 1977 sous le nom de GEPAN, puis rebaptisé SEPRA, celui-ci est revenu en 2005 à son acronyme précédent, auquel a été ajouté le « I » de « information ».

Les enquêtes menées par le GEIPAN ont permis d'élucider un certain nombre de cas, mais selon ses dirigeants passés et présents, des centaines d'autres restent inexplicables. Faut-il en conclure que des objets d'une nature inconnue sillonnent notre atmosphère ?

On peut d'abord se demander si les travaux du GEIPAN sont solides. Des phénomènes extraordinaires requièrent des preuves extraordinaires. Or, entre l'absence de recherche exhaustive des sources de méprise possibles, le « syndrome du témoin parfait » ou les délais d'intervention importants, l'analyse des rapports d'enquête du GEIPAN laisse entrevoir des lacunes méthodologiques récurrentes.

Revenons à l'appel de novembre dernier. Dans le groupe de personnalités figurait un ancien pilote de ligne d'*Air France*, témoin d'un étrange phénomène. Le 28 janvier 1994, l'équipage du vol AF-3532 observe au-dessus de Paris ce qui lui paraît être un objet gigantesque. Or, au moment de l'observation, un radar enregistre un écho non identifié. Le GEIPAN conclut à la présence d'un engin d'une technologie inconnue, la piste radar venant confirmer l'observation visuelle.

Un examen des faits montre qu'en réalité l'écho non identifié se situe à droite et à proximité de l'*Airbus*, alors que l'ovni est observé au loin sur sa gauche. Les deux phénomènes ne sont donc pas corrélés. Les caractéristiques de l'écho – fuyatif, faible et non confirmé – suggèrent un aéronef dépourvu de système d'identification radar, en phase d'atterrissage.

De plus, l'enquête menée ne permet pas d'exclure la possibilité que les témoins oculaires aient pu prendre pour un ovni un objet banal – avion ou ballon-sonde – vu dans des conditions particulières.

On retrouve ici l'une des plus fréquentes lacunes qui entachent les travaux du GEIPAN : les sources de méprise envisageables pour un cas précis ne sont quasiment jamais vérifiées de façon exhaustive. Conscient de la fragilité du témoignage humain, le GEIPAN avait pourtant établi dès sa création une méthodologie stricte, afin d'extraire des données fiables des rapports d'observation. Par exemple, il s'interdisait d'enquêter sur des cas ne présentant qu'un témoignage unique dépourvu de toute espèce de confirmation matérielle. Mais devant l'extrême rareté de ce type de preuve, le mieux est devenu l'ennemi du bien et les enquêteurs ont souvent mis en avant des traces physiques dont, à l'instar de l'écho radar du vol AF-3532, le lien causal avec l'ovni allégué était fragile ou inexistant.

Les témoins sont rarement remis en question

Un effet physique aussi contestable sera mis en avant pour étayer le cas dit de « l'Amarante » en 1982. Un biologiste nancéen prétendait qu'un ovni avait endommagé certaines plantes de son jardin (des amarantes). L'altération en question, étayée par des échantillons mal conservés, ressemble à une simple déshydratation après un coup de gel. Seule la qualité de scientifique du témoin a empêché le GEIPAN de mettre en cause la crédibilité de son récit, dont maints détails suggèrent pourtant une origine purement subjective – illusion visuelle complexe, voire hallucination.

Le même « syndrome du témoin parfait » réapparaît dans l'affaire du vol AF-3532. La profession du témoin est utilisée pour écarter d'autorité une possibilité de méprise, comme si les pilotes d'aéronefs étaient immunisés contre les illusions perceptives et cognitives. Cette idée conduit fréquemment les enquêteurs à ne pas rechercher d'autres témoins susceptibles de confirmer ou d'infirmer les dires du témoin « qualifié », voire d'identifier ce qu'il a observé.

Le délai d'intervention du service – en théorie inférieur à 48 heures – atteint souvent plusieurs semaines ou plusieurs mois. L'enquête sur le cas AF-3532 n'a été menée

que trois ans après sa médiatisation. Pour la légendaire « rencontre du troisième type » de Cussac, les enquêteurs se rendront sur les lieux 11 ans après l'événement ! Un tel laps de temps a un impact inévitable sur les témoignages recueillis, tels l'apparition de faux souvenirs ou l'altération des vrais.

Les conclusions des enquêtes ne sont par ailleurs presque jamais remises en cause, même lorsque surviennent des éléments nouveaux susceptibles d'expliquer l'affaire. Il en est ainsi du cas de Nort-sur-Erdre en 1987, où un jeune garçon avait enregistré sur une cassette audio le bruit de l'ovni qu'il disait avoir vu. Bien qu'il ait admis depuis 2005 avoir tout inventé, les responsables du GEIPAN présentent toujours en 2007 le cas comme « inexplicable ».

En définitive, malgré les bonnes intentions qui ont présidé à sa naissance, le GEIPAN n'a que trop rarement respecté la démarche rigoureuse dont il s'est toujours réclamé. De fait, les enquêtes présentées comme les plus probantes, largement médiatisées, souffrent toutes de l'un ou l'autre des biais méthodologiques précédemment décrits.

Une véritable admission par le GEIPAN de ses errements passés se traduirait par une requalification de nombreux cas prétendus « inexplicables ». Ce travail permettrait une mise à jour bénéfique de la méthodologie forgée dans les premières années d'existence du service. L'emploi de méthodes standards, scientifiquement validées, pour contrôler les estimations physiques livrées par les témoins oculaires (tailles angulaires, durées, etc.) représenterait par exemple un grand pas en avant. Resterait ensuite aux enquêteurs actuels à appliquer cette méthodologie renouée avec la rigueur qui a souvent fait défaut à leurs devanciers. Dès lors, le GEIPAN serait en mesure d'effectuer correctement la première partie de sa mission : enquêter sur les phénomènes aérospatiaux non identifiés.

Qu'en est-il de la seconde, l'information ? Le site du GEIPAN, qui offre depuis avril 2007 aux internautes l'accès aux archives des enquêtes du service, ne renseigne toujours pas réellement le public sur les différents phénomènes susceptibles d'être pris pour des ovnis. Il n'informe pas non plus sur les enquêtes récentes. Les témoins, les médias et le public peuvent ainsi toujours entretenir la croyance en un mystère pourtant déjà élucidé par les enquêteurs. Le problème ne peut qu'empirer avec l'augmentation du nombre de rapports attendus par le service, qui incite désormais les pilotes et astronomes amateurs à lui transmettre toute observation de phénomène aérospatial non identifié.

Il est maintenant temps que le GEIPAN évite les travers du passé, donne la priorité au qualitatif sur le quantitatif et assume une information publique cohérente, transparente et suivie. Tels sont du moins les vœux que l'on peut formuler pour qu'il ait un avenir, utile au grand public comme aux scientifiques.

David ROSSONI est archiviste et diplômé en histoire, **Éric MAILLOT** est collaborateur du Laboratoire de zététique de l'Université de Nice et **Éric DÉGUILLAUME**, diplômé en histoire des sciences, est membre de l'Association L'Observatoire zététique.

D. ROSSONI, E. MAILLOT et E. DÉGUILLAUME, *Les OVNI du CNES : 30 ans d'études officielles (1977-2007)*, éditions Book-e-Book.com, décembre 2007.

Site internet du GEIPAN : www.cnes.fr/geipan/

■ POUR LA SCIENCE

et le Palais de la Découverte

vous invitent
le mardi 19 février 2008
à 18 h 30

Les limites de la lutte contre le tabagisme

Avec Marion Adler

**Responsable de l'Unité de consultation
de tabacologie de l'Hôpital
Antoine Béchère, à Clamart**

Au Palais de la Découverte
Av. Franklin-D.-Roosevelt
75008 Paris
www.palais-decouverte.fr

Entrée libre
Inscription obligatoire à
tabac@palais-decouverte.fr

En partenariat avec



Kelvin, Perry et l'âge de la Terre

Si les scientifiques avaient davantage pris en considération l'un des critiques contemporains de Kelvin, la théorie de la dérive des continents aurait pu s'imposer quelques décennies plus tôt.

Les scientifiques du XIX^e siècle se sont longuement divisés sur la question de l'âge de la Terre, qui n'a reçu de réponse définitive qu'au milieu du XX^e siècle. La plus célèbre estimation de l'époque victorienne est due au renommé physicien William Thomson (1824-1907), connu après 1892 sous le nom de lord Kelvin. Elle est notoirement fautive : 24 à 400 millions d'années.

L'histoire de Kelvin et de l'âge de la Terre est souvent racontée comme un combat à la David contre Goliath, des géologues jouant le rôle du plus faible, armé seulement de la frêle épée du raisonnement géologique, tandis que lord Kelvin leur assénait de grands coups de physique mathématique, auréolée de prestige. Nombreux sont ceux qui pensent que le calcul de Kelvin était erroné parce qu'il ne connaissait pas la radioactivité. Nous montrerons que l'erreur n'était pas là. La faille dans le raisonnement de Kelvin a été identifiée en 1894 par l'un de ses propres assistants, John Perry, dont les idées auraient pu faire considérablement progresser l'étude de la Terre... si seulement les géologues de l'époque les avaient comprises et appréciées à leur juste valeur.

Dans son traité *Théorie analytique de la chaleur*, le mathématicien français Joseph Fourier (1768-1830) avait posé les bases quantitatives de l'estimation de l'âge de la Terre par Kelvin. Ce dernier s'est attaqué pour la première fois à l'âge de la Terre en 1844, en montrant que la mesure du taux de déperdition de chaleur à la surface de la planète pouvait imposer des limites à son âge.

Kelvin a imaginé que la Terre s'était solidifiée à partir d'un état initialement fondu, de telle sorte que ses conditions initiales étaient celles d'une température élevée uniforme partout, avec une surface se maintenant par la suite à une température constante. Avec ces hypothèses, la température dépend de la profondeur sous terre et du temps écoulé depuis cet état initial.

Pour se faire une idée des lois physiques en jeu, imaginez que vous sortiez une dinde de Noël brûlante du four et que vous la placiez immédiatement au congélateur (un congélateur parfait, capable d'extraire la chaleur instantanément et qui reste toujours à sa température de réglage). Initialement, la dinde serait partout à la même température (à supposer que vous l'ayez rôtie assez longtemps). Seule une mince

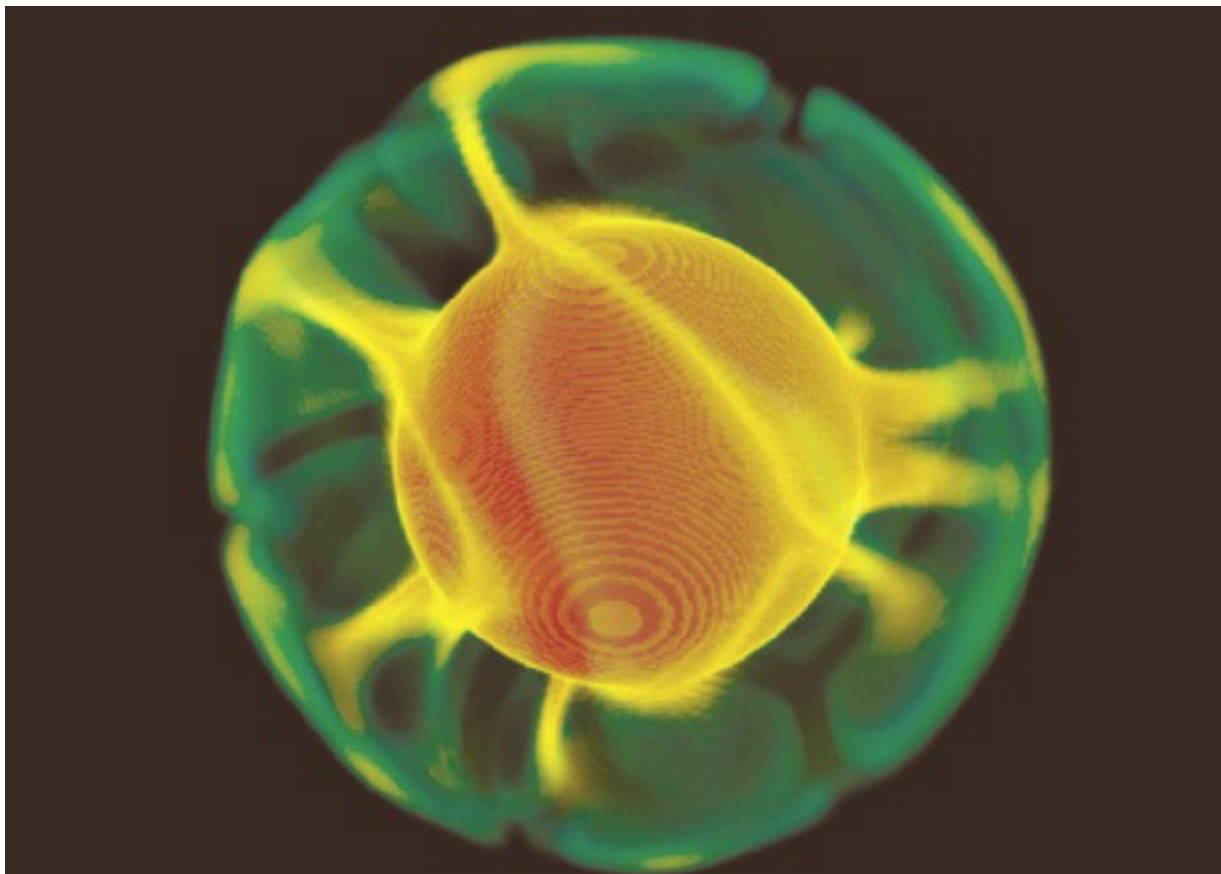
couche superficielle se retrouverait tout de suite à la température du congélateur. Mais rapidement, les couches externes de la dinde se refroidiraient à mesure que la chaleur diffuse vers l'extérieur, bien que le centre se maintienne à sa température initiale, identique à celle du four. À la fin, bien sûr, tout finirait par se refroidir, y compris le cœur de la farce ; en d'autres termes, la température à l'intérieur de la dinde dépend à la fois de la distance à la surface et du temps écoulé depuis que la volaille a été placée dans le congélateur.

L'analyse de Kelvin permet de chiffrer cette expérience de pensée. Fourier avait montré que les changements de température dans un solide obéissent à l'équation de diffusion : la température en un point donné varie à une vitesse proportionnelle à la dérivée spatiale seconde (la dérivée de la dérivée) de la température ; la constante de proportionnalité, qui dépend du matériau, est le coefficient de diffusion thermique.

Quand une dinde rôtie refroidit...

Les solutions de l'équation de diffusion montrent que le temps nécessaire pour que la chaleur se propage sur une distance donnée est proportionnel au carré de cette distance divisé par le coefficient de diffusion thermique. Par exemple, cinq minutes après avoir placé le volatile dans votre congélateur, seule une couche de viande épaisse d'environ un centimètre aura refroidi ; tout ce qui se trouve plus en profondeur reste à la température de rôtissage. Mais alors qu'il faut cinq minutes pour que la chaleur s'évacue par conduction thermique du premier centimètre de la dinde, il en faut 20 pour extraire la chaleur des deux premiers centimètres (c'est pourquoi de nombreux livres de cuisine recommandent de sortir la dinde du four jusqu'à une heure avant de la découper, pour que l'intérieur continue à cuire sans que l'extérieur se dessèche).

Supposons que la différence de température entre l'intérieur du four et l'intérieur du congélateur soit de 200 degrés : 180 °C est une température correcte pour rôtir une dinde, et - 20 °C convient pour conserver des glaces. Cinq minutes après le début de l'expérience, le centimètre extérieur de la dinde s'est refroidi, et le gradient de température avec la profondeur est de 20 degrés par millimètre de dinde. Au bout de 20 minutes, deux centimètres se sont refroidis, et le gra-



1. À l'échelle d'une vie humaine, le manteau terrestre est aussi rigide que de l'acier, mais sur des millions d'années, il se comporte comme un fluide très visqueux, qui transporte la chaleur de l'intérieur vers la surface. Les détails de ce mouvement restent à déterminer, mais la simulation numérique illustrée ici suggère que des panaches de roche chaude s'élèvent du fond du manteau, et que du matériau plus froid coule vers le

noyau brûlant de fer liquide (*en rouge*). L'avènement de la théorie des plaques tectoniques, dans les années 1960, a permis de comprendre que les continents se déplacent en réponse à ces mouvements de convection dans le manteau. Comme le décrivent les auteurs, John Perry, un ancien assistant de Kelvin, a proposé en 1895 que ce même phénomène expliquerait l'écart entre l'estimation de l'âge terrestre par Kelvin et celle des géologues.

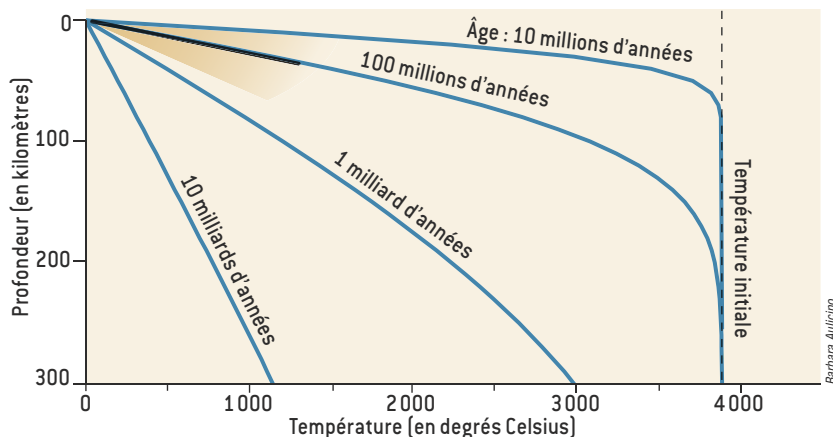
dient est de dix degrés par millimètre de dinde. Une autre façon d'exprimer cette relation est de dire que le gradient thermique est inversement proportionnel à la racine carrée du temps que la dinde a passé dans le congélateur.

En inversant ce type de calcul, Kelvin a pu obtenir l'âge de la Terre en fonction du gradient géothermique à la surface de la planète. Pour filer l'analogie, si vous découvriez dans votre congélateur une dinde de Noël, vous pourriez déterminer depuis combien de temps elle y est en mesurant le gradient de température à sa surface. Un gradient de température de cinq degrés par millimètre, par exemple, impliquerait un « âge » de 80 minutes (de même, les amateurs de romans policiers savent qu'en mesurant la température d'un cadavre, on peut déterminer l'heure du décès).

Lorsque Kelvin a présenté pour la première fois de tels arguments en 1844 et 1846, il ne disposait pas de mesures

fiables du gradient géothermique. Mais quand il s'est à nouveau penché sur le problème 15 ans plus tard, des gradients géothermiques avaient été mesurés en plusieurs endroits dans le monde. Kelvin choisit alors pour ses calculs un gradient moyen de 1/50° de degré Fahrenheit par pied, soit environ 36 °C par kilomètre. Il estima la température initiale de la Terre à 3 900 °C, à partir d'expériences de fusion de roches, et des expériences de laboratoire lui fournirent la diffusivité thermique de matériaux typiques de la croûte terrestre. Compte tenu des incertitudes sur les valeurs du gradient géothermique et du coefficient de diffusion thermique, Kelvin obtint pour la Terre un âge compris entre 24 et 400 millions d'années.

Les scientifiques accordent à leurs conclusions une plus grande confiance s'ils y parviennent par plusieurs chemins indépendants. C'était le cas pour Kelvin, qui avait également considéré l'âge du Soleil. Selon les connaissances



Barbara Aulicino

2. Selon le raisonnement de William Thomson, plus connu sous le nom de lord Kelvin, le gradient géothermique (le taux d'élévation de la température avec la profondeur) dépend du temps écoulé depuis que la planète était une boule uniformément chaude de roche fondue. Peu après que la Terre s'est solidifiée, l'écorce externe a commencé à se refroidir, engendrant un fort gradient thermique près de la surface. Avec le temps, la chaleur diffuse vers l'extérieur, et le

gradient thermique diminue progressivement. Ainsi, dans le modèle de Kelvin, les mesures du gradient géothermique près de la surface (*région du graphique soulignée de marron*) permettent de déterminer l'âge de la Terre. Kelvin avait initialement calculé que le gradient géothermique moyen (*souligné d'un trait dans le haut du graphique*) correspond à un âge d'environ 100 millions d'années. Il a ultérieurement revu son estimation à la baisse, à environ 20 millions d'années.

de l'époque, l'énergie rayonnée par le Soleil ne pouvait provenir que de l'énergie potentielle de gravitation libérée au cours de l'accrétion de cette étoile. Kelvin avait calculé cette énergie et conclu que le Soleil ne pouvait maintenir son flux actuel de rayonnement pendant plus de 100 millions d'années, un chiffre en accord avec l'âge de la Terre et obtenu de façon indépendante. Ainsi, bien qu'il ait par la suite revu à la baisse son estimation, à 20 millions d'années environ, Kelvin resta convaincu que la Terre était âgée de quelques dizaines ou centaines de millions d'années, pas plus.

Or les géologues savent aujourd'hui que la Terre a environ 4,5 milliards d'années. Où Kelvin s'est-il trompé ? A-t-il utilisé des valeurs erronées pour le gradient géothermique, pour la diffusivité thermique des roches ou pour la température initiale de la Terre ? Rien de tout cela. Si l'on refaisait les calculs de Kelvin avec les valeurs actuelles, plus précises, de ces paramètres, on obtiendrait encore un âge compris entre 24 et 96 millions d'années.

L'idée d'un âge fini s'impose

Avant d'examiner les arguments de Kelvin, il est utile de décrire la vision du monde à laquelle il s'opposait. Les géologues du début du XIX^e siècle pensaient généralement que la Terre était d'âge illimité. Cette doctrine permettait aux géologues d'expliquer tout phénomène non par les lois de la physique, mais par ce que le géologue et éducateur américain Thomas Chrowder Chamberlin qualifiait en 1899 de « retraits inconsidérés à la banque du temps ». Pour Kelvin, ce jeu sans règles n'était pas de la science et était en contradiction avec les lois de la thermodynamique, qu'il avait largement contribué à dégager.

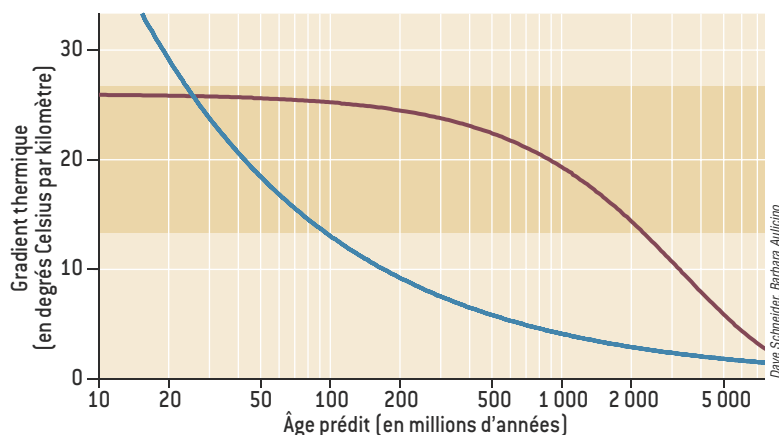
La géologie a connu d'énormes progrès du simple fait de devoir affronter Kelvin au sujet de l'âge de la Terre. À la fin du XIX^e siècle, les géologues ont fini par admettre que l'âge de la Terre était fini et que l'estimation de sa valeur par un raisonnement quantitatif était une des clefs de la démarche géologique. Mais personne d'autre que Kelvin ne s'était attelé à un tel raisonnement avant 1895.

Un principe unique sous-tend tous les arguments de Kelvin concernant l'âge de la Terre : la conservation de l'énergie. Pour ses analyses, Kelvin ajoutait trois hypothèses, dont deux ne s'appliquaient qu'à ses arguments sur la Terre : la rigidité de la planète et l'homogénéité de ses propriétés physiques. La troisième hypothèse, à savoir qu'il n'y avait pas de source inconnue d'énergie, s'appliquait à la fois à la Terre et au Soleil. Nous savons aujourd'hui que cette troisième hypothèse explique l'erreur de Kelvin sur l'âge du Soleil : l'énergie rayonnée par le Soleil provient de la fusion thermonucléaire de l'hydrogène en hélium au sein de l'étoile.

L'histoire officielle veut que la troisième hypothèse de Kelvin ait également faussé ses calculs de l'âge de la Terre. S'il est vrai que la désintégration d'éléments radioactifs dans les profondeurs terrestres constitue une source persistante de chaleur, ce n'est pas l'ignorance de cette source qui a conduit à une estimation erronée de l'âge de la Terre. La véritable faille de l'argument de Kelvin fut montrée du doigt en 1894 par l'un de ses anciens assistants, John Perry, presque dix ans avant que la radioactivité ne soit reconnue comme une source de chaleur.

Irlandais du Nord, Perry fit ses études à l'Université de Belfast où il assista aux cours d'ingénierie du frère de Kelvin, James Thomson, qui était lui aussi un scientifique de renom. Après l'obtention de son diplôme, Perry passa quatre années en tant qu'instituteur avant de devenir l'assistant de Kelvin à Glasgow. Il garda toute sa vie beaucoup d'affection et de respect envers Kelvin, ne serait-ce qu'en raison de son influence sur sa carrière. En moins d'une année, l'instituteur était devenu professeur d'ingénierie à Tokyo, où il contribua grandement à la naissance de l'industrie japonaise, avant de revenir à Londres occuper une chaire de professeur.

Perry était un infatigable enseignant d'ingénierie et de mathématiques, et insistait pour que les mathématiques soient présentées de manière claire et concrète. Des générations d'élèves ont été, sans le savoir, influencées par son approche, parce qu'il reconnut l'énorme potentiel du papier millimétré pour noter les données et tracer les fonctions, et



Science Museum Pictorial

3. John Perry, professeur d'ingénierie et ancien assistant de Kelvin, est parvenu à concilier les mesures du gradient géothermique avec un âge terrestre de plusieurs milliards d'années, en considérant que la planète a un manteau fluide convectif recouvert d'une couche relativement mince de roche solide. Les calculs utilisant le modèle de Perry (*courbe bordeaux*) avec une croûte solide de 50 kilomètres et intégrant les données actuelles sur la diffusi-

vité thermique et le point de solidification des roches du manteau montrent que la fourchette d'estimations du gradient géothermique moyen (*bande marron sur le graphique*) est compatible avec des âges allant jusqu'à deux ou trois milliards d'années. Avec le modèle de conductivité thermique uniforme de Kelvin (*courbe bleue*), les gradients géothermiques ne sont compatibles qu'avec des âges compris entre une vingtaine et une centaine de millions d'années.

en en faisant baisser le prix de sa production, il a généralisé son utilisation. Il acheva sa carrière en tant que professeur d'ingénierie au *Royal College of Science* (l'actuel *Imperial College* de Londres).

Perry, rechignant à mettre dans l'embarras son patron, essaya d'abord de convaincre Kelvin par le dialogue et la correspondance privée. Mais Kelvin le repoussa, soit qu'il ne comprenait pas l'approche de Perry, soit qu'il n'était pas intéressé. Perry décida donc en 1895 de rendre publique l'affaire et écrivit à la revue *Nature*. Au lieu de se focaliser sur les calculs de Kelvin, Perry suggérait plutôt de revenir sur ses hypothèses.

Convection dans les profondeurs

Dans le modèle de Kelvin, le gradient thermique près de la surface de la Terre diminue à mesure que la couche externe refroidie s'épaissit avec le temps. Si la Terre avait beaucoup plus que 100 millions d'années, cette couche serait tellement épaisse que le gradient thermique serait bien plus faible que la valeur observée. Perry se rendit compte qu'il y avait un moyen simple de ne pas épaissir la couche externe : il proposa que la chaleur pouvait être transférée beaucoup plus efficacement à l'intérieur de la Terre qu'à la surface. Si tel était le cas, l'intérieur profond fournirait d'importantes réserves de chaleur, ce qui maintiendrait longtemps en surface un gradient thermique élevé, et l'estimation de l'âge de la Terre par Kelvin serait alors trop basse, peut-être d'un facteur important.

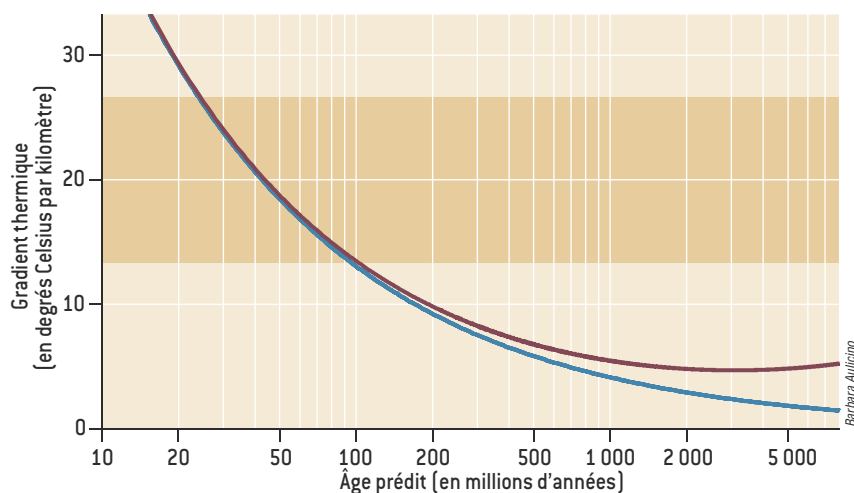
Perry proposa plusieurs mécanismes possibles pour ces transferts internes de chaleur. En particulier, la convection (déplacements de matière) dans l'intérieur fluide, ou partiellement fluide, de la Terre, pouvait transférer la chaleur beaucoup plus efficacement que la diffusion thermique. Perry fit remarquer que si seule une mince couche superficielle de notre planète était solide, et que le reste de la Terre était un fluide convectif, alors l'intérieur serait bien mélangé et aurait partout la même température. Le modèle de Perry remplace ainsi la dinde de notre expérience virtuelle par une

bouteille de même gabarit que la dinde contenant, par exemple, du cidre chaud. Seuls les quelques millimètres externes de cet objet sont solides (le verre de la bouteille), et la convection brasse l'intérieur. Ainsi, le gradient de température près de la surface peut rester élevé pendant une longue durée.

Le calcul de Perry montre que si la Terre a une croûte solide de 50 kilomètres d'épaisseur, soumise à la conduction thermique, surmontant un manteau fluide parfaitement convectif, alors les gradients thermiques mesurés près de la surface sont compatibles avec un âge de deux à trois milliards d'années. Reconnaissant que le transfert de chaleur dans le manteau ne pouvait être parfaitement efficace, Perry modélisa par la suite l'intérieur profond comme un solide doté d'une « quasi-diffusivité » élevée. Ses résultats étaient en accord avec le calcul simple initial, qui suggérait que la Terre pouvait être âgée de plusieurs milliards d'années. Des calculs complets de la convection dans le manteau (impossibles avant l'avènement des ordinateurs) confirment que le raisonnement de Perry était valide, même s'il ne peut conduire à des âges aussi précis que ceux déterminés aujourd'hui sur la base des radiodatations.

En résumé, Perry était parvenu à concilier un calcul physique de l'évolution thermique de la Terre avec le grand âge que les géologues préconisaient. Pour cela, Perry avait juste eu besoin d'introduire l'idée que la chaleur circulait plus facilement dans les profondeurs de la Terre que dans ses couches externes.

Pourtant, jusqu'à aujourd'hui, la plupart des géologues restent persuadés que l'erreur (compréhensible) de Kelvin était de ne pas avoir pris en compte la radioactivité interne de la Terre. En 1903, Pierre Curie et son assistant Albert Laborde déterminaient que le processus de désintégration radioactive nouvellement découvert libérait de la chaleur, et on ne tarda pas à affirmer que cette source d'énergie suffisait à bousculer la conclusion de Kelvin sur l'âge de la Terre. En particulier, Ernest Rutherford, considéré comme le père de la théorie nucléaire, s'exprima sur la question lors d'une rencontre de la *Royal Institution* en 1904 :



Ramsay, Muspratt, AIP Emilio Segre Visual Archives, Brittle Books Collection

4. En 1904, Ernest Rutherford suggéra que les estimations par Kelvin de l'âge terrestre étaient trop basses parce qu'il ignorait l'existence de la chaleur dégagée par radioactivité. Certes, les connaissances actuelles sur la production de chaleur par radioactivité dans le man-

teau terrestre modifient la prédiction du modèle de Kelvin (*courbe bleue sur le graphique*), mais très peu (*courbe bordeaux*) : des âges supérieurs à environ 100 millions d'années restent incompatibles avec les gradients géothermiques mesurés (*bande marron*).

« J'entrai dans la pièce, qui était plongée dans une semi-obscurité, et remarquai bientôt lord Kelvin dans l'assistance ; je me rendis compte que j'allais avoir quelques difficultés avec la dernière partie de mon intervention concernant l'âge de la Terre, où mes vues étaient en contradiction avec les siennes. À mon grand soulagement, il s'endormit profondément, mais au moment où j'abordais le point important, je vis le vieil oiseau se redresser sur son siège, ouvrir un œil et me lancer un regard torve ! Une inspiration me vint alors, et je déclarai que lord Kelvin avait « limité l'âge de la Terre à condition qu'aucune source nouvelle de chaleur ne soit découverte. Cette formulation prophétique désigne justement ce que nous examinons ce soir, le radium ! » Et là, merveille ! Le visage du vieux bonhomme s'épanouit en un large sourire. »

Peut-être la célébrité de cette anecdote explique-t-elle l'idée toujours répandue que la découverte de la chaleur radioactive sapait à la fois une hypothèse de base du calcul de Kelvin et sa conclusion. Mais cette affirmation est incorrecte. La découverte de la chaleur radioactive n'aurait invalidé la conclusion de Kelvin que si l'incorporation de cette chaleur dans son calcul devait produire un âge nettement différent pour la Terre.

La radioactivité ne résout rien

On estime aujourd'hui la chaleur produite par la radioactivité interne à environ 2×10^{13} watts, ce qui équivaut à la moitié environ de la quantité totale de chaleur qui s'échappe de la planète à l'heure actuelle. Il semblerait donc que la production de chaleur interne de la Terre puisse contribuer notablement aux gradients de température près de la surface. Il faut cependant se rappeler que cette production de chaleur est distribuée dans tout le volume des 2 900 kilomètres d'épaisseur du manteau, et que le calcul de Kelvin montrait que la diffusion ne peut extraire en 100 millions d'années que la chaleur des 100 premiers kilomètres de la Terre. Si l'on refaisait les calculs de Kelvin en incluant la radioactivité, on devrait par conséquent s'attendre à ce que seule la chaleur créée dans les 100 kilomètres externes de la Terre (environ cinq

pour cent du volume du manteau) contribue au gradient de température de la surface. Ainsi, même si Kelvin avait inclus la chaleur radioactive dans ses calculs, son estimation de l'âge de la Terre ne serait pratiquement pas affectée. Un calcul détaillé prenant en compte la contribution de la radioactivité confirme cet argument simple.

L'analyse par Perry de l'âge de la Terre suivait l'exemple de Kelvin en cela qu'il insistait sur l'analyse d'un système physique simple, mais le système qu'il avait choisi d'analyser différait de celui de Kelvin. Un manteau fluide était la condition nécessaire (largement reconnue) de l'isostasie : l'idée (datant du XIX^e siècle) que les éléments de la croûte tels que les chaînes de montagnes doivent flotter sur le manteau plus dense, de la même façon que les icebergs flottent sur l'eau. Et pourtant, Kelvin était sûr de lui-même en rejetant la notion d'un intérieur fluide : il savait, grâce à l'étude des marées, que les 1 500 premiers kilomètres au moins du manteau étaient aussi rigides que l'acier. Perry s'évertua à faire changer Kelvin d'avis sur ce point, avec un langage spécifiquement choisi pour le toucher. Il écrivit :

« [...] la base réelle de votre calcul est votre hypothèse qu'une Terre solide ne peut changer de forme [...] même en 1 000 millions d'années, sous l'action de forces ayant tendance à modifier sa forme, et pourtant nous observons la fermeture progressive des passages dans les mines, et nous savons que des plissements, des failles et d'autres changements se produisent constamment dans la Terre solide sous l'action de ces forces persistantes. Je sais que la roche solide n'est pas comme de la poix de cordonnier, mais 10^9 années, c'est une longue durée, et les forces sont grandes ! »

La référence de Perry à la poix de cordonnier était délibérée et mérite une explication. Kelvin, comme beaucoup de physiciens de l'époque, pensait que la lumière ne pouvait pas se propager dans le vide, mais nécessitait un milieu physique, l'« éther ». Pour permettre aux ondes lumineuses de se propager, ce milieu était censé avoir des propriétés élastiques aux très petites échelles de temps, mais il devait être mou aux échelles de temps plus grandes, afin que la Terre puisse s'y déplacer librement.

Sans pouvoir trouver de formulation mathématique satisfaisante pour l'éther, Kelvin aimait recourir à une démonstration physique des propriétés requises. Il plaçait de l'eau dans un cylindre de verre, faisait flotter quelques petits bouchons de liège à la surface, les recouvrait d'une couche de poix de cordonnier écossaise, et enfin posait des balles de fusil tout en haut. Sur une courte période d'observation, il ne se passait rien de visible, mais au bout de six mois les bouchons et les balles s'étaient enfouis dans la poix. Après un an, on pouvait trouver les bouchons en haut, et les balles en bas. Ainsi, la poix présente de la rigidité à court terme, mais elle est molle à long terme.

Ces propriétés étaient exactement celles dont Kelvin avait besoin pour l'éther. Perry soutenait que le désaccord entre Kelvin et les géologues pouvait être résolu si l'on supposait le manteau terrestre rigide aux courtes échelles de temps, mais fluide sur des périodes plus longues. L'argument échappa complètement à Kelvin et, semble-t-il, à tous les autres.

Perry publia ses arguments dans *Nature*, revue qui jouait en 1895 un rôle aussi important qu'aujourd'hui dans la communication des scientifiques. D'autres acteurs du débat sur l'âge de la Terre en ont donc eu certainement connaissance. Pourquoi alors l'hypothèse des mouvements de convection au sein de la Terre, avancée par Perry, n'a-t-elle pas été admise dans la décennie qui précéda la découverte de la chaleur radioactive, ni de fait ultérieurement ?

Perry, trop respectueux ?

Des facteurs personnels ont probablement joué un rôle. La plupart des scientifiques modernes, s'ils étaient aussi convaincus que Perry d'avoir raison dans un débat majeur, auraient davantage persévéré. Mais Perry révérait Kelvin et il n'était pas du genre à se mettre en avant dans un débat. Il semble qu'il ait laissé tomber la question et qu'il ait passé les 25 années suivantes à enseigner ou faire des recherches sur des sujets très variés, avant de succomber au scorbut contracté lors d'un long voyage en mer... entrepris pour des raisons de santé.

Une telle déférence pour les grands noms ne se rencontre plus guère, mais Perry n'était pas le seul à l'époque. C'est peut-être une atmosphère de déférence comparable qui a permis à la vision de Rutherford de supplanter celle de Kelvin après 1904.

L'indifférence face à l'argument de Perry s'explique peut-être également par le fait que cet argument n'était pas compris de la plupart de ceux qui s'intéressaient à l'âge de la Terre. Les biologistes se sont emparés de la conclusion que l'estimation de Kelvin pouvait être fausse, tout en admettant qu'ils ne comprenaient pas le raisonnement de Perry. Les géologues se sont contentés de considérer la démonstration que Kelvin se trompait comme une preuve que, pour ce qui concerne la Terre « réelle », l'intuition géologique l'emporte sur la simple analyse physique (mais n'oublions pas que les géologues de l'époque ont produit des estimations de l'âge de la Terre tout aussi fausses, voire plus, que celle de Kelvin).

De nombreux géologues avaient le sentiment que les physiciens étaient incapables de comprendre le raisonnement géologique et que la géologie est trop complexe pour être réduite à un modèle mathématique. L'analyse de Perry

montre pourtant que l'erreur de Kelvin n'était pas d'appliquer une physique simple, mais d'appliquer un modèle physique simple qui était inadéquat.

Kelvin comme Perry savaient qu'en science, les modèles simples – destinés à analyser les phénomènes dans leurs grandes lignes, et non dans leurs moindres détails – sont des outils indispensables. Ce que Kelvin n'a pas reconnu, c'est que tous les modèles simples sont voués à s'effondrer, à des degrés divers, et que les scientifiques ont autant à apprendre de leurs échecs que de leurs réussites.

Un mythe aidé par une anecdote

L'analyse de Perry montrait que l'échec du modèle de conduction thermique pour la Terre s'expliquait si le manteau était le siège de mouvements de convection. Si la communauté scientifique de l'époque avait intégré le message de Perry, les premières radiodatations de la Terre auraient confirmé l'explication convective du gradient géothermique observé. La prise de conscience qu'il y a de la convection à l'intérieur d'une Terre solide en apparence aurait mis en difficulté la vision « fixiste » de la Terre, qui soutenait qu'il était physiquement impossible que les continents se déplacent horizontalement. La conception fixiste a considérablement freiné le progrès de la géologie durant la première moitié du XX^e siècle : les partisans de la dérive des continents ont dû sans cesse argumenter, face à un considérable scepticisme, que la Terre solide pouvait s'écouler sur de très longues périodes. Il fallut attendre les années 1960 pour que leurs vues s'imposent.

Pourquoi tant de scientifiques, y compris des géologues, continuent de croire que la découverte de la radioactivité a simultanément apporté la preuve de l'erreur de Kelvin et l'explication de cette erreur ? La réponse est peut-être en partie que cela fait une belle histoire. Selon Arthur Eve, l'un des premiers biographes de Rutherford, ce physicien influent aimait à répéter l'épisode où il avait dû faire preuve de présence d'esprit face à ce « vieil oiseau » de Kelvin. Il est tout à fait possible que la forme plaisante de l'anecdote, et la renommée de son auteur, aient conduit à une acceptation peu critique du mythe. Il est difficile d'empêcher les scientifiques vieillissants, qui développent un goût pour les anecdotes, de répéter à l'envi des histoires qu'ils trouvent amusantes. Mais leurs jeunes collègues doivent se garder de confondre histoires et vérité historique.

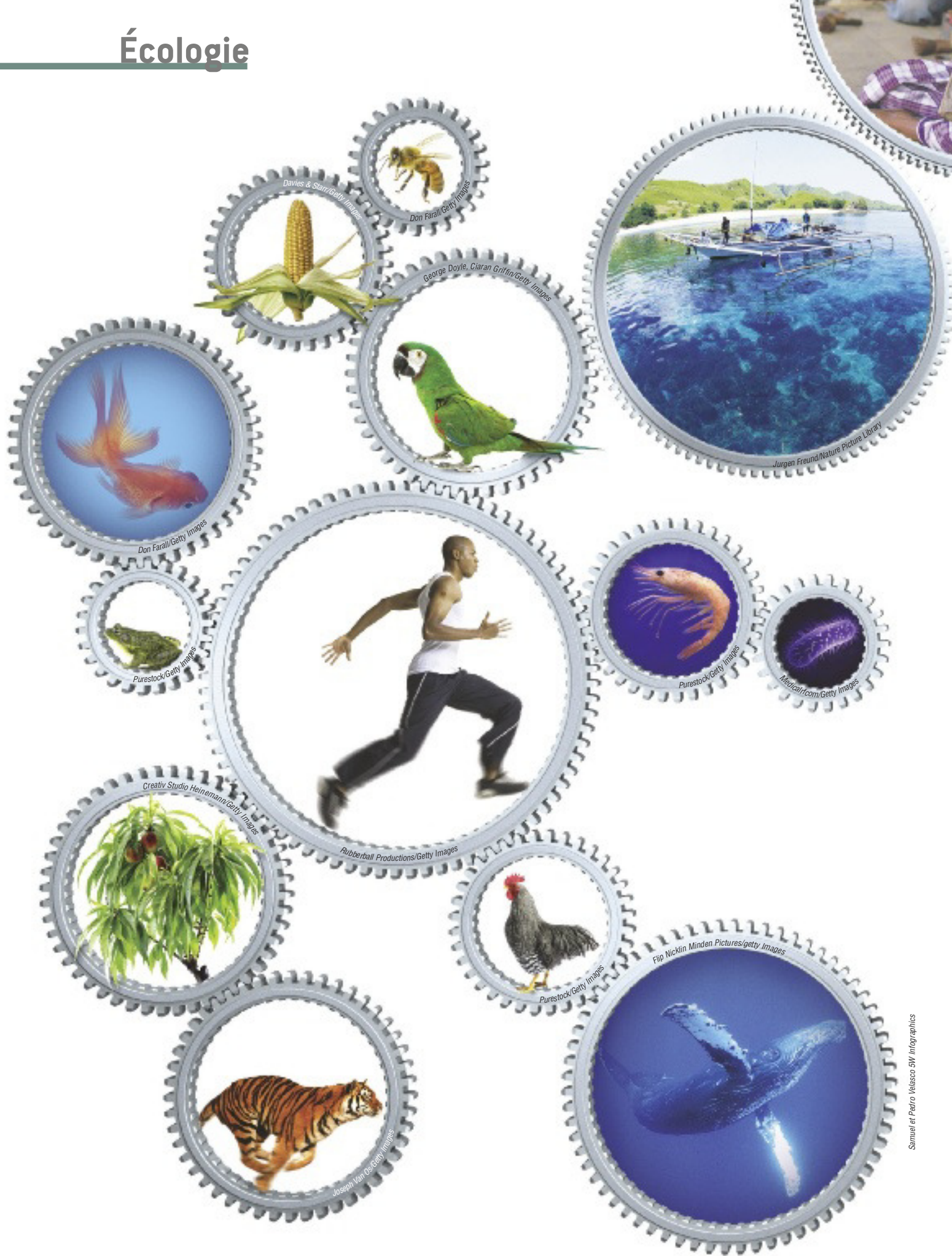
Nous remercions la revue *American Scientist* de nous avoir autorisés à publier cet article.

Philip ENGLAND est professeur au Département de géologie de l'Université d'Oxford, en Grande-Bretagne. **Peter MOLNAR** est professeur au Département de sciences géologiques de l'Université du Colorado, à Boulder, aux États-Unis. **Frank RICHTER** est professeur au Département de sciences géophysiques de l'Université de Chicago, aux États-Unis.

P. C. ENGLAND, P. MOLNAR et F. M. RICHTER, *John Perry's neglected critique of Kelvin's age for the Earth : A missed opportunity in geodynamics*, in *GSA Today*, vol. 17(1), pp. 4-9, 2007.

D. LINDLEY, *Degrees Kelvin*, Joseph Henry Press, 2004.

C. L. E. LEWIS et S. J. KNELL (éd.), *The age of the Earth : from 4004 BC to AD 2002*, The Geological Society, 2001.





Repenser l'écologie

*Il est illusoire, voire contre-productif, de préserver
un écosystème isolé du reste du monde.
Replaçant l'homme au centre de leur stratégie,
les écologues repensent les interactions
de tous les acteurs de la biodiversité.*

Peter Kareiva • Michelle Marvier

La *World Conservation Union* est le réseau international de conservation des espèces le plus vaste du monde : il regroupe 83 États et des centaines d'agences, gouvernementales ou non. En 2004, cette organisation a inscrit trois vautours – le vautour à long bec, le vautour à bec fin, le vautour oriental à dos blanc – sur la liste des espèces menacées. La population totale de ces trois vautours atteignait presque 40 millions d'individus en Inde et en Asie du Sud au début des années 1990 ; elle a aujourd'hui chuté de plus de 97 pour cent. Pourquoi doit-on protéger les vautours de l'extinction ? Parce que c'est une obligation éthique que de sauvegarder la biodiversité de la planète. Mais d'autres raisons plus originales peuvent être invoquées...

On a longtemps ignoré ce qui provoquait le déclin de la population de ces vautours. Certains observateurs pensaient que la réduction de l'habitat ou la pollution en étaient responsables. Mais, il y a quelques années, on a découvert que ces animaux étaient tués par un médicament anti-inflammatoire couramment administré aux vaches. Chez les bovins et chez l'homme, ce médicament réduit la douleur ; chez les vautours, il provoque une insuffisance rénale.

Comme les vautours ont quasiment disparu, les centaines de milliers de carcasses de bovins habituellement « nettoyées » par ces rapaces ont pourri au soleil et ont été consommées par des chiens. La population de chiens sauvages a donc explosé – et avec elle, le risque de transmission de la rage. Ainsi, le destin de millions de personnes est lié à celui des vautours ; sauver les vautours de l'extinction protégerait l'homme d'une redoutable maladie.

Il n'est pas évident de voir le lien entre le bien-être de l'homme et la protection des espèces menacées, mais ce genre de relations existe dans de nombreuses situations

où l'environnement est en danger. Des écosystèmes tels que les marais et les mangroves protègent les hommes des tempêtes ; les forêts et les récifs coralliens sont des sources de nourriture (gibier, poissons) et de revenus. Un dommage causé à un écosystème particulier peut nuire à une autre partie du monde, même éloignée, aussi bien qu'aux personnes dont les ressources dépendent de cet écosystème.

Malgré ces interdépendances, le public et certains gouvernements considèrent de plus en plus que les efforts pour préserver la biodiversité des végétaux et des animaux l'emportent sur tout le reste, et notamment sur les besoins de l'homme. Afin d'inverser cette tendance – et de mieux servir l'humanité –, un nombre croissant de partisans de la protection de l'environnement (dont nous faisons partie) soutient que les anciennes méthodes définissant les priorités devraient être abandonnées, au profit d'une approche mettant l'accent sur la sauvegarde d'écosystèmes précieux pour l'homme. Notre programme devrait permettre de sauver de nombreuses espèces, tout en protégeant la santé et les moyens de subsistance des hommes.

Des points très chauds...

Pourquoi ceux qui veulent protéger l'environnement ont-ils une réputation de « misanthropes » ? Parce que des millions de personnes ont dû quitter leurs terres ou se sont vu privées de leurs sources de nourriture et de revenus afin que des animaux et leur habitat soient protégés. La décision controversée du président du Kenya, Mwai Kibaki, de restituer le parc Amboseli à ses habitants originaux, les Massaïs (des éleveurs semi-nomades), entraîne un mécontentement croissant vis-à-vis des déplacements de ces populations. En Asie et en Afrique, les chasseurs et les agriculteurs soutiennent que ces parcs naturels

limitent leurs ressources alimentaires et leurs revenus. Les fermiers et les bûcherons américains sont en colère, car ils craignent de perdre leurs privilèges pour l'eau ou leur emploi à cause des saumons ou des hiboux...

Cette notion de « nature contre homme » est issue d'une stratégie de protection de l'environnement concentrée sur ce que l'on nomme les « points chauds ». En 1988, Norman Myers, de l'Université d'Oxford, a développé l'idée de points chauds de biodiversité, de petites zones renfermant une grande variété d'espèces végétales endémiques, c'est-à-dire caractéristiques d'une région peu étendue. N. Myers a utilisé la diversité de ces végétaux comme instrument de mesure, car les classifications des plantes sont fiables et représentent souvent les seules données disponibles. De surcroît, la diversité des plantes refléterait celle des animaux. Ainsi, N. Myers et ses collègues ont identifié 25 points chauds – la région de Cerrado au Brésil en est une ; la Corne de l'Afrique en est une autre – sur lesquels doivent se concentrer les projets de préservation de l'environnement.

Auparavant, les campagnes de protection étaient centrées sur des espèces animales emblématiques, tels les pandas, les baleines et les phoques. À l'inverse, le concept de points chauds offre un ensemble de critères rigoureux et quantifiables qui permettent de définir les priorités pour la protection de l'environnement. Ce système de triage, plus scientifique que des photographies d'animaux émouvants, est fondé sur le décompte des espèces. Cette nouvelle approche semble aussi plus réaliste : les organismes de protection de la nature disposent de fonds limités et investissent désormais là où la plupart des espèces peuvent être sauvegardées. Durant ces 15 dernières années, cette stratégie a été adoptée par les organisations écologiques locales et internationales.

Toutefois, David Orme, de l'Imperial College à Londres, a récemment fait remarquer que les « points chauds » seraient de mauvais indicateurs : des zones regorgeant

d'espèces végétales n'abritent pas nécessairement de nombreuses espèces animales. Marcel Cardillo, aussi à l'Imperial College, a noté que les animaux vivant dans les points chauds riches en fleurs ne sont pas les plus menacés.

D'autres biologistes ont montré que, parmi les régions du globe présentant le moins de diversité végétale, beaucoup offrent des habitats saisonniers, des « escalas » sur les routes migratoires ou des sites de nidification. Ainsi, un demi-million de manchots de Magellan se rassemblent tous les mois de septembre dans la Punta Tomba, en Argentine. Cette région de Patagonie, aride et plantée d'arbustes, possède peu de plantes endémiques, et ne mériterait même pas le qualificatif de « point tiède » ! Pourtant l'économie locale a besoin des manchots qui viennent y nicher : 70 000 touristes les admirent chaque année.

Il existe de nombreuses régions semblables – des sites de faible biodiversité végétale – qui sont néanmoins cruciales pour certaines espèces animales importantes d'un point de vue écologique ou économique : des étendues de toundra qui accueillent des canards, des cygnes et des oies ; des rivières tempérées où vivent des saumons.

Évaluer les services rendus par les écosystèmes

D'autres principes sont nécessaires pour protéger l'environnement. Bien qu'on ne comprenne pas toujours le concept de biodiversité, on tient à la nature en tant que source de nourriture, de combustible, de matériaux de construction, de loisir et d'inspiration. Les scientifiques ont quantifié ce capital naturel selon le principe des « services écologiques » (ou *Ecosystem services* en anglais).

Quels services l'écosystème rend-il à l'homme ? Il fournit des produits pour lesquels il existe des marchés, tels les plantes médicinales et le bois de construction, et assure des mécanismes dont l'importance économique est souvent sous-estimée : la filtration des eaux, la pollinisation, la régulation du climat, le contrôle des inondations et des maladies, et la formation des sols. Quand Robert Costanza, de l'Université du Vermont, et d'autres économistes ont essayé d'attribuer une « valeur » à tous ces processus naturels, ils ont découvert que la valeur marchande annuelle de ces « services écologiques » dépassait le produit intérieur brut de tous les pays réunis !

Et les scientifiques ne sont pas les seuls à s'intéresser aux services écologiques ; de plus en plus de gouvernements et d'organisations non gouvernementales considèrent la sauvegarde de ces services comme un objectif fondamental. En 2000, les Nations unies ont commandé une étude sur les services écologiques. Un an après, une équipe internationale de plus de 1 300 scientifiques a entrepris l'une des missions « écologiques » les plus ambitieuses : le *Millennium Ecosystem Assessment*. Ce projet a montré l'impact que les hommes ont eu sur les services écologiques au cours des 50 dernières années.

Les services sont classés en quatre catégories : les services d'approvisionnement (la fourniture de produits tels que la nourriture), les services de régulation (des fonctions de régulation tel le contrôle des inondations), les ser-

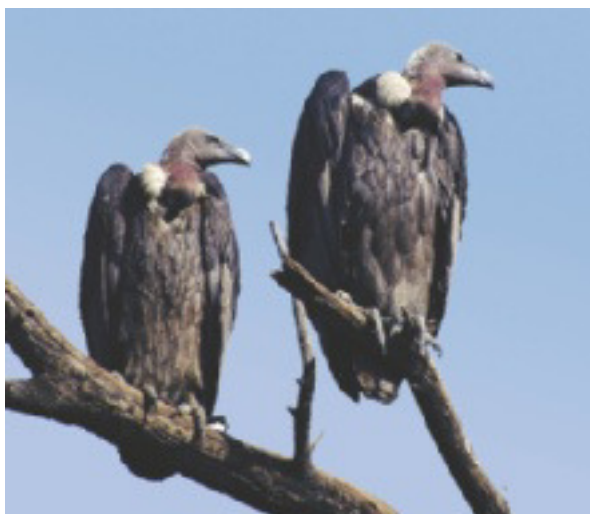


1. Les pays indiqués en rouge recèlent des écosystèmes qui, selon les auteurs, sont des priorités pour la protection de l'environnement. La conservation et la restauration de ces écosystèmes amélioreraient les conditions de vie des habitants. Les auteurs ont identifié ces sites en utilisant des données sur la pauvreté, l'importance des ressources naturelles pour l'économie locale et la dégradation des sites.

vices culturels (l'apport d'avantages non matériels, telle une sensation de bien-être) et les services de soutien (qui apportent les éléments de base de l'écosystème, telle la formation des sols). Cette étude montre que non seulement la plupart des services écologiques ont décliné, mais aussi qu'ils sont utilisés de façon non appropriée.

Pour le public, le tsunami qui a sévi dans l'océan Indien en 2004 et l'ouragan Katrina en 2005 ont mis en lumière les relations entre les écosystèmes et les conditions de vie des hommes. Dans les deux cas, les dommages ont été amplifiés par la disparition de la végétation naturelle. La destruction, depuis 70 ans, d'environ 4 000 kilomètres carrés de marécages en Louisiane a exacerbé la tempête engendrée par Katrina. En Asie du Sud-Est, la reconversion des forêts de la mangrove côtière en réserves de crevettes a supprimé la barrière faisant office de brise-lames contre les tsunamis. Et le Sri-Lankais Farid Dahdouh-Guebas, de l'Université Vrije à Bruxelles, a montré que les côtes ayant des forêts de mangrove intactes n'avaient subi presque aucun dégât.

Ni les marécages de Louisiane ni les mangroves du Sri Lanka ne comptaient parmi les points chauds de la biodiversité du monde : ils renferment très peu d'espèces de



Ashok Jan/Nature Picture Library

2. Quand les vautours d'Inde commencèrent à disparaître, personnes n'avait envisagé que cela aurait des conséquences pour l'homme. Mais les vautours ne nettoyaient plus les carcasses animales, de sorte que la population des chiens sauvages a explosé et, avec elle, le risque de transmission de la rage. On avait oublié que préserver la vie sauvage, c'est aussi protéger l'homme.

Protéger l'environnement : histoire de deux stratégies

La stratégie de protection de la nature consistant à se concentrer sur les « points chauds » (à gauche) néglige les écosystèmes précieux pour la santé et le développement de l'espèce humaine. L'approche des « services écologiques » (à droite) permettrait de sauvegarder l'environnement en tenant compte de la présence de l'homme, et définirait de nouvelles priorités.

La stratégie des « points chauds »



Tut De Roy/Minden Pictures

L'idée fondamentale

Elle consiste à identifier les régions menacées qui présentent une grande diversité végétale et à les protéger, en supposant qu'agir ainsi profite aussi aux populations d'animaux plus difficiles à répertorier que les plantes. À ce jour, 25 régions qualifiées de points chauds ont été identifiées, dont le parc Bocaina, au Brésil (ci-dessus).

Comment s'y prend-on ?

On crée un parc ou une réserve naturelle pour protéger la vie animale et végétale. On décourage les gens de vivre dans cette zone ou de l'exploiter. La zone est surveillée et les frontières renforcées.

Les inconvénients

Les régions riches en espèces végétales ne sont pas forcément riches en espèces animales. Les populations locales sont souvent déplacées ou perdent d'importantes ressources. Cette stratégie n'est pas soutenue par les populations locales.

La stratégie des « services écologiques »



Gitan-Jean

L'idée fondamentale

On insiste sur la dépendance des populations vis-à-vis d'écosystèmes variés – comme c'est le cas avec les revenus du tourisme dans la Punta Tomba, en Argentine (ci-dessus) – et on identifie les écosystèmes qui sont menacés et dont la dégradation nuit aux habitants de ces régions.

Comment s'y prend-on ?

Là où les écosystèmes sont en train de se dégrader, on établit un plan de conservation protégeant l'écosystème et profitant à la communauté humaine qui dépend de cet écosystème.

Les avantages

Comme le public sait qu'il dépend de différents écosystèmes pour sa santé et sa sécurité, il sera plus enclin à soutenir les projets de protection de l'environnement. La biodiversité sera donc préservée, mais pas aux dépens des hommes.

Protéger les communautés et les habitats pauvres

Le problème

Les marais salants, les lits d'algues et les récifs à huîtres de la côte du golfe de Floride abritent des lamantins, des tortues de mer et d'autres espèces menacées, tout comme ils servent de « nurserie » aux crevettes, aux crabes et aux vivaneaux (des poissons), d'une grande importance économique. Ces habitats protègent aussi contre les tempêtes qui accompagnent les ouragans, tel Dennis en 2004 (voir la photographie page ci-contre). Pourtant, les stratégies visant à défendre et à réhabiliter ces écosystèmes côtiers – qui seraient aussi utiles aux personnes et permettraient d'étendre l'habitat d'espèces menacées ou précieuses pour l'économie – ont souvent été ignorées, au profit de projets industriels qui ont accéléré l'érosion et la disparition de l'habitat.

La solution

Des scientifiques des organisations *Nature Conservancy* et *National Oceanic and Atmospheric Administration* ont combiné une carte des marais côtiers en situation critique et une carte des espèces menacées dans une zone côtière de Floride, avec une carte de prévision des tempêtes et une carte où sont concentrées les communautés pauvres les plus exposées à ces tempêtes (voir ci-dessous).

En superposant ces différentes données, ils ont identifié les régions dont la réhabilitation protégerait les habitants les plus pauvres en même temps que de nombreuses espèces parmi les plus importantes de ces zones.



plantes endémiques, et on estime que le nombre d'espèces végétales et animales qu'ils contiennent n'atteint pas le dixième de celui trouvé dans les forêts humides.

Pourtant, la disparition d'un habitat peut avoir des conséquences économiques importantes. Les vents qui soufflent sur le Sahara et le Sahel – ces déserts d'Afrique qui s'étendent constamment – transportent des poussières vers l'Est, par-dessus l'océan Atlantique. Chaque année, plusieurs centaines de millions de tonnes de ce sable se déposent sur les Amériques et les Antilles. Les poussières, les polluants, les micro-organismes et les nutriments accompagnant ce sable participent à la destruction des récifs coralliens – ce qui perturbe l'industrie du tourisme et de la pêche. Un pâturage excessif et de mauvaises pratiques agricoles en Afrique septentrionale et subsaharienne y ont aggravé la pauvreté, la famine et la malnutrition et ont détérioré les coraux et l'économie de la moitié du monde.

Soulignons que les bienfaits économiques des services écologiques sont avant tout nécessaires aux nations en développement. En effet, ces pays tirent une grande partie de leurs revenus du bois, des fibres et de l'agriculture; la sylviculture et la pêche sont cinq à dix fois plus importantes en tant que composantes de l'économie de ces nations que pour les États-Unis et l'Europe. Un rapport des Nations unies montre que la protection de l'environnement est fondamentale pour diminuer la pauvreté de 750 millions de personnes vivant dans les zones rurales du monde.

La santé de l'homme est aussi en danger quand les écosystèmes et les cycles naturels sont perturbés. Les vautours d'Inde mentionnés en introduction ne sont qu'un exemple parmi des centaines. Presque deux millions de personnes meurent chaque année à cause d'un approvisionnement en eau insuffisant. La sauvegarde des marais et des forêts réduirait le nombre de ces décès: les marais sont des filtres natu-

rels qui améliorent la qualité de l'eau pour l'alimentation et l'agriculture; les forêts retiennent une partie des sédiments qui troublent l'eau. En outre, la protection des forêts et des prairies réduirait les nuages de poussière venus d'Afrique et ceux – encore plus gros – qui traversent l'océan Pacifique depuis l'Ouest de la Chine; ces derniers seraient une cause de l'augmentation des cas d'asthme aux États-Unis.

Voyons un autre exemple plus évident du lien qui existe entre la dégradation de l'écosystème et la santé de l'homme: intéressons-nous aux micro-organismes qui deviennent pathogènes pour l'homme, alors qu'ils ne le colonisaient pas auparavant. Les deux tiers des maladies émergentes dans le monde, telles la fièvre Ebola et la grippe aviaire, sont dus à des organismes qui infectent des animaux et qui entrent en contact avec l'homme à cause de modifications de l'utilisation des sols et des pratiques agricoles. Et il ne s'agit pas que de maladies « exotiques »... En exterminant les loups et les pumas, les habitants de l'Est des États-Unis ont déclenché une explosion de la population des daims et des tiques des daims, ce qui a engendré plus de 20 000 nouveaux cas de la maladie de Lyme par an. L'éradication de ces prédateurs il y a plus d'un siècle met aujourd'hui en danger la santé de l'homme.

Notre approche de la protection de l'environnement fondée sur les services écologiques représente un nouvel « emballage » des idées classiques. Mais elle diffère par plusieurs aspects. Tout d'abord, nous pensons que de nombreuses personnes qui protègent la nature refusent de voir le monde tel qu'il est et qu'elles doivent cesser de s'accrocher à leur vision d'une nature sauvage déconnectée de tout. De plus en plus de forêts sont défrichées et de marais asséchés pour développer l'agriculture. De plus en plus d'espèces marines sont pêchées, ce qui entraîne leur diminution. La biodiversité décline. La nature « sauvage » et l'homme sont intimement imbriqués.

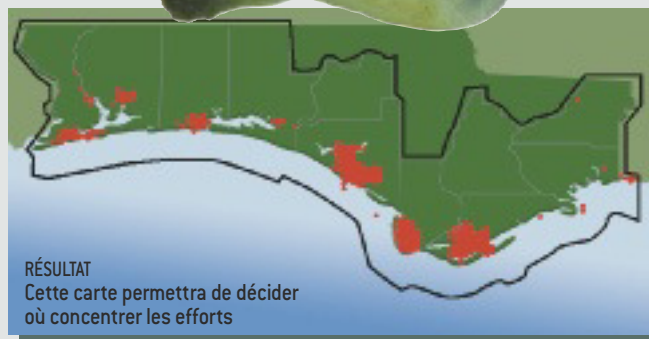


Marl Darr Welch Ap Photo



Le résultat

En montrant comment faire coïncider la protection de la nature et les besoins des hommes, les scientifiques commencent à recueillir une forte adhésion du public aux efforts de protection et de réhabilitation de cette zone côtière de Floride.



Étant donné que notre environnement est constitué de systèmes soumis à l'influence de l'homme, on doit envisager la protection de la biodiversité dans le contexte de paysages incluant des centres urbains, une agriculture intensive, des forêts et des rivières gérées par l'homme, et pas simplement dans le contexte de la seule conservation des espèces. Aujourd'hui, ce sont les régions protégées qui nécessitent la gestion la plus intensive si l'on veut conserver leur « état sauvage ». Ainsi, dans le parc Kruger, en Afrique du Sud, un parc bien géré, on a été obligé de réduire la taille des populations d'éléphants pour éviter le surpeuplement.

La nouvelle stratégie : concilier l'homme et la nature

Par ailleurs, il faudrait que la protection de l'environnement soit focalisée sur les régions où la dégradation des écosystèmes menace le plus le bien-être de l'homme, tels les mangroves en Asie, les marais dans le Sud-Est des États-Unis, les zones arides de l'Afrique subsaharienne et les récifs coralliens un peu partout dans le monde. Cette approche serait très efficace là où les organismes gouvernementaux et les associations de protection de l'environnement cherchent à travailler ensemble, à la fois pour la protection des populations et pour la préservation de l'environnement.

Ainsi, en Floride, deux associations ont pour objectif d'identifier des régions qui présentent un intérêt commun pour le bien-être de la population et la protection de l'environnement. En dressant une carte des habitats selon leur capacité à protéger les communautés humaines en plus de leur biodiversité, les participants à ce programme ont découvert d'importantes régions à préserver.

Tous ceux qui veulent protéger la nature devraient collaborer plus étroitement avec les experts en développement. Depuis 20 ans, de nombreux projets de développement durable ont cherché à concilier les intérêts de ces deux groupes ; mais ils se sont seulement intéressés aux produits déjà commercialisés, tels des poissons ou certaines ressources forestières, et ont négligé tous les autres services écologiques. En réunissant et en coordonnant l'énergie et les capitaux des différents partenaires, on renforcerait l'efficacité et l'impact de ces deux approches. Par exemple, les investissements pour purifier l'eau sont souvent les mêmes que ceux susceptibles de protéger la biodiversité aquatique.

Qui plus est, si l'on néglige les problèmes sociaux, les mesures de protection de la biodiversité ont peu de chances de rencontrer l'adhésion du public. Les écologues doivent aller au-delà de leur tendance à placer l'environnement dans un vase clos et à le tenir éloigné de toute autre préoccupation – faute de quoi, ils sont condamnés à l'impuissance.

Enfin, les tentatives de protection de l'environnement que nous envisageons seront évaluées non seulement par le nombre d'espèces protégées, mais aussi par l'amélioration du bien-être des hommes. De telles évaluations ont déjà commencé. En 1980, le gouvernement indonésien et l'organisation *Nature Conservancy* ont déclaré le parc Komodo réserve naturelle, non seulement pour protéger le dragon de Komodo, une espèce menacée, mais aussi pour préserver les forêts et les récifs coralliens. La recette des entrées dans le parc est consacrée à des projets de développement local et à la recherche de nouvelles sources de revenus : la culture d'algues, le tourisme, la sculpture sur bois et l'élevage de poissons de récifs. En 2006, une enquête auprès des villageois habitant près de Komodo a montré que la majorité d'entre eux soutiennent cette zone protégée à cause des nouveaux revenus qu'elle a engendrés.

Protéger la faune

Le problème

En Namibie, les populations San, souvent nommées Bushmen, sont marginalisées et très pauvres. Les enfants souffrent souvent de retard de croissance. Les San ont été déplacés de leurs habitats traditionnels et, privés de moyens durables de gagner leur vie, ont été contraints de braconner et de pratiquer une chasse excessive. Le rhinocéros noir, l'une des espèces les plus menacées du monde, a été une des victimes de cette situation.

La solution

En 1996, le gouvernement namibien a voté une loi qui donne aux San la propriété du gibier et de tous les revenus provenant du tourisme et des produits de la chasse. Des programmes de protection, couvrant 17 pour cent du territoire namibien et incluant 60 communautés, afin de gérer la faune – par exemple en repérant les déplacements du gibier – ont été instaurés. Et on a donné des fonds pour permettre aux San de participer à ces actions locales de protection.

Le résultat

Dans les régions où ces programmes de protection ont été bien exécutés, les populations d'éléphants, de zèbres, d'oryx et d'antilopes ont été multipliées par six. La Namibie abrite maintenant la plus grande population de rhinocéros noirs au monde. En outre, plus de 500 emplois à plein temps et plus de 3 000 emplois à temps partiel ont été créés. En 2004, le tourisme et la chasse ont rapporté 2,5 millions de dollars de revenus. Ce cas illustre aussi certains des défis auxquels se heurte l'approche des « services écologiques » pour la protection de l'environnement : de nombreux San restent marginalisés, et certains observateurs pensent que ces populations indigènes devraient devenir propriétaires du territoire, et pas seulement de la faune.

Pourquoi faut-il préserver les espèces rares ? Parce qu'elles jouent un rôle crucial : elles sont une forme d'assurance. À cause du bouleversement climatique global et des profondes modifications des sols, les espèces rares d'aujourd'hui peuvent devenir les espèces abondantes de demain ; c'est pourquoi on doit en sauver autant que possible. En Californie, l'abeille européenne (non indigène) est le plus important pollinisateur d'un point de vue économique. Si la population d'abeilles européennes devait décroître considérablement (et elle a été récemment menacée par l'introduction de tiques), certaines populations d'abeilles indigènes, moins abondantes, pourraient augmenter et remplir le rôle essentiel de pollinisateurs.

Quelles priorités pour quel environnement ?

Bien qu'il soit inacceptable que les êtres humains laissent s'éteindre toutes les espèces exceptées celles qui rendent, ou pourraient rendre, des services, il est irréaliste de penser qu'on peut ramener le monde à un état préindustriel. Certaines extinctions dues à l'homme sont inévitables, et on doit être réaliste sur ce qu'on peut et ce qu'on ne peut pas faire. Il faut d'abord protéger les écosystèmes là où la biodiversité offre des services aux personnes qui en ont vraiment besoin.

Nous pensons qu'au lieu de se focaliser sur les 10 ou 25 régions à protéger parce qu'elles sont riches en plantes indigènes, ceux qui protègent l'environnement devraient plutôt identifier les écosystèmes « bouées de sauvetage », c'est-à-dire les régions très pauvres où l'économie dépend des systèmes naturels, mais où les ressources de l'écosystème sont dégradées. Les efforts visant à approvisionner les populations en eau potable, à réduire l'érosion des sols et à éviter la surpêche aideront les populations et protégeront la biodiversité. Par ailleurs, certains organismes devraient

Protéger l'eau potable

Le problème

Une grande partie de l'eau alimentant Quito, la capitale de l'Équateur, vient des montagnes andines qui abritent une immense variété de plantes et d'animaux endémiques, notamment le condor des Andes. Bien qu'une réserve de condors ait été créée (voir la photographie), elle est mal gérée. En aval, de nombreuses régions autour de la ville n'ont pas suffisamment d'eau pour satisfaire leurs besoins, et, dans la ville, l'eau n'est pas partout potable. Les mauvaises méthodes de culture et d'exploitation forestière au voisinage de la réserve de condors et le pâturage d'animaux de ferme trop près du fleuve en sont responsables.

La solution

En 2000, plusieurs partenaires ont créé un programme pour améliorer la distribution de l'eau potable. Une société d'hydroélectricité, le fournisseur d'eau de la ville de Quito, et une taxe de deux pour cent payée par les habitants de Quito ont financé ce programme. On a ainsi collecté 4,9 millions de dollars pour soutenir la protection de l'environnement, l'éducation et des projets d'épuration d'eau en amont de Quito.

Le résultat

Cette année, 11 nouveaux gardes ont été embauchés pour patrouiller dans la zone protégée et on a entrepris un grand programme d'éducation pour enseigner aux agriculteurs de meilleures méthodes d'utilisation des sols. On a planté plus de 3,5 millions d'arbres afin de reboiser les bassins hydrographiques. Il est encore trop tôt pour savoir si ces mesures amélioreront la qualité de l'eau, mais on est en train de créer un réseau de stations de contrôle de l'eau, et les populations locales se sentent très concernées.



Pete Oxford/Nature Picture Library

continuer à promouvoir la conservation d'espèces et de régions sans utilité apparente... En effet, insister sur les services qu'un écosystème peut rendre à l'homme ne signifie pas qu'il faille radicalement changer les objectifs de la protection de l'environnement ; il faut juste renforcer le soutien du public et revoir les priorités.

Reste à montrer que les actions visant à protéger les écosystèmes peuvent favoriser le développement économique. L'avenir de notre approche en tant que stratégie de protection dépend de la collaboration – improbable ! – entre les écologistes et les experts financiers. Toutefois, le monde des affaires manifeste beaucoup d'enthousiasme pour cette approche. En novembre 2005, par exemple, le groupe *Goldman Sachs* a annoncé un programme écologique qui prévoit la mise à disposition d'une somme de un milliard de dollars pour investir dans les énergies renouvelables, évaluer l'impact de ces projets sur les services écologiques et mettre en place un comité d'experts pour prospecter des marchés écologiques.

La Banque mondiale encourage aussi les nations à adopter des méthodes de comptabilité « écologique », où l'évaluation de l'économie et de la productivité nationale inclut des mesures qui créditent les services écologiques et pénalisent la dégradation engendrée par la pollution par exemple. Cette évaluation économique et cette création de marchés pour les services écologiques fournissent un système de données quantifiables auquel les entreprises et les particuliers peuvent se référer ; ce qui en soi est une amélioration par rapport aux approches fondées sur des espèces animales emblématiques ou sur des espèces végétales endémiques.

Certaines voix – comme celles du prix Nobel de la paix en 2004, Wangari Muta Maathai, et de l'ancien secrétaire général des Nations unies Kofi Annan – ont attiré l'attention sur la relation entre l'environnement, la prospérité de l'humanité et la paix dans le monde. Kofi Annan a déclaré que « notre combat contre la pauvreté, l'inégalité et la maladie est directement lié à la santé de la planète elle-même ». Les partisans de la protection de l'environnement doivent entendre ce message et le relayer. La protection de la nature ne sera vraiment internationale et soutenue que lorsque l'homme sera au centre de sa mission.

Pour la Science et l'émission *Science publique* s'associent pour vous proposer chaque mois un débat en direct sur *France Culture*. Ce débat reprendra le thème d'un article publié dans *Pour la Science*.

Ainsi, Michel Alberganti vous proposera en février :
Peut-on préserver la biodiversité contre l'homme ?

Voir informations complémentaires sur le site de *Pour la Science*.

Peter KAREIVA est directeur scientifique de l'organisation *Nature Conservancy* et **Michelle MARVIER** est professeur à l'Université de Santa Clara, en Californie, où elle dirige l'Institut des études environnementales.

R. BARBAULT, *Espaces protégés : stratégie dépassée ou dispositifs d'avenir ?*, in *Pour la Science*, vol. 359, pp. 14-17, septembre 2007.

How much is an ecosystem worth ? Assessing the economic value of conservation, World Bank, 2005.

Valuing ecosystem services : toward better environmental decision making, National Research Council. National Academic Press, 2005.

M. PALMER et al., *Ecology for a crowded planet*, in *Science*, vol. 304, pp. 1251-1252, 2004.

les **mardis** conférences grand public
de l'Institut Curie

Mardi 19 février 2008

**Vers de nouvelles thérapies
plus ciblées contre le cancer**



**Cycle de conférences
le dernier mardi
de chaque mois
de 18 à 20 heures
à l'Institut Curie
Amphithéâtre Constant-Burg
12 rue Lhomond
75005 Paris**

- **Dr Laurent Mignot,**
chef du département
de Médecine oncologique.
- **Dr Patricia De Crémoux,**
responsable de l'unité
de Pharmacogénétique.
- **Dr Jean-Yves Pierga,**
médecin-chercheur, responsable
de l'Hôpital de jour.

Les thérapies ciblées s'attaquent
directement aux défaillances
des cellules tumorales. Quels sont
les changements apportés
par ces molécules dans la pratique
quotidienne des cancérologues ?
Et quel est leur avenir ?

Infos : 01 44 32 40 64
Tout le programme sur www.curie.fr





Les sursauts gamma, témoins de

L'observation détaillée des sursauts gamma immédiatement après leur détection a récemment permis de découvrir les objets qui en sont à l'origine, et de préciser les mécanismes qui produisent ces bouffées de rayons gamma visibles depuis le fond de l'Univers.

Robert Mochkovitch

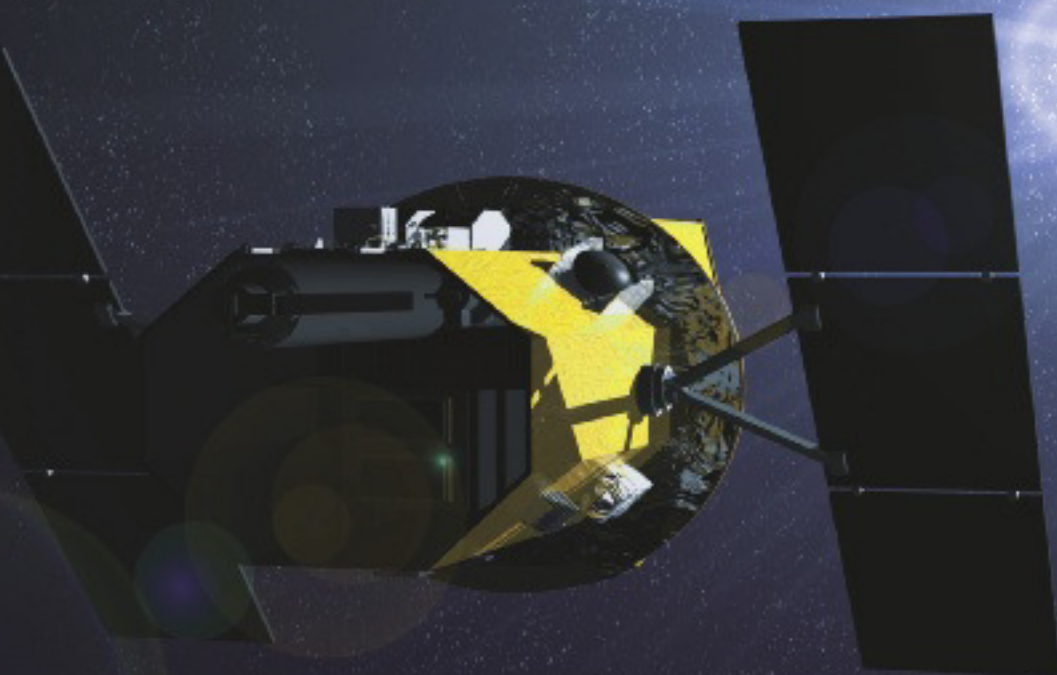
Quelque part dans l'espace, deux étoiles à neutrons, résidus d'étoiles massives ayant explosé, gravitent l'une autour de l'autre depuis des milliards d'années. Au fil du temps, perdant de l'énergie, elles se rapprochent insensiblement. Jusqu'au jour où, trop proches, elles tombent brusquement l'une sur l'autre et fusionnent dans une explosion cataclysmique, qui brille pendant une fraction de seconde d'un éclat un milliard de milliards de fois plus puissant que celui du Soleil. Telle est peut-être l'origine des sursauts gamma courts. Qu'ils soient courts (quelques dixièmes de seconde) ou longs (plusieurs dizaines de secondes), les sursauts gamma, des explosions visibles dans la partie la plus énergétique du spectre lumineux, intriguent les scientifiques depuis leur découverte il y a 40 ans.

Ces dernières années, les progrès en matière d'observation ont permis d'observer les sursauts gamma dans toute une gamme de longueurs d'onde quelques minutes à peine après leur détection. Ces observations ont forcé les théoriciens à adapter les modèles décrivant les sursauts, et suggèrent une origine distincte pour les sursauts longs et courts. Nous allons étudier dans cet article ces avancées récentes.

Dès le début, l'étude des sursauts gamma a été pleine de surprises et de rebondissements. Tout commence comme un roman d'espionnage : en 1963, l'URSS, les États-Unis et la Grande-Bretagne signent un traité interdisant les essais nucléaires dans l'atmosphère. Ne faisant pas confiance aux Soviétiques, les Américains lancent peu après un programme de satellites espions appelés *Vela*, capables de détecter les essais nucléaires *via* la bouffée de rayonnement gamma qui accompagne l'explosion. Et de fait, le 2 juillet 1967, le satellite *Vela 4a* détecte une bouffée de photons gamma... Pourtant, aucune bombe soviétique n'avait explosé ! D'abord tenue secrète par les militaires, la découverte de ce « sursaut gamma » ne fut rendue publique que six ans plus tard. Entre-temps, 15 autres sursauts avaient été détectés, qui avaient permis d'exclure que ces explosions aient lieu dans l'atmosphère terrestre ni même dans le Système solaire.

Très vite, des programmes d'observation sont mis en place depuis l'espace (les photons gamma ne peuvent traverser l'atmosphère terrestre). Ils permettent de dégager les principales caractéristiques de ces mystérieux événements. Les sursauts gammas durent de quelques millisecondes à quelques centaines de secondes. Les sursauts durant plus de deux secondes sont qualifiés de sursauts longs, les autres, de sursauts courts. Leurs profils temporels sont variés, comprenant

la mort violente des étoiles



parfois un unique pic de luminosité, parfois une succession de pics enchevêtrés. La distribution en énergie des photons émis (le spectre) suit en général une loi de puissance présentant une cassure (la pente est plus prononcée dans les hautes énergies). L'énergie caractéristique des photons émis lors d'un sursaut, représentée par la position de cette cassure dans le spectre, est de l'ordre de quelques centaines de kiloélectronvolts, dans le domaine gamma du spectre électromagnétique.

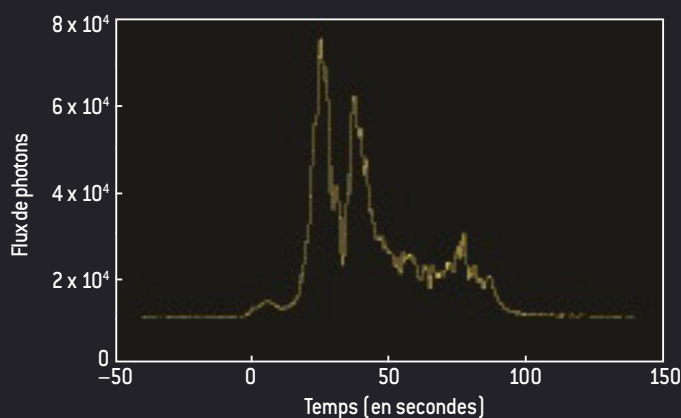
La révolution des rémanences

À partir de ces indices, les théoriciens échafaudent les premiers modèles. La brièveté des sursauts, la rapidité de leurs variations – de l'ordre de la milliseconde – et l'émission dans le domaine gamma, tout suggérait une source compacte, telle une étoile à neutrons (un astre qui rassemble près de 1,5 fois la masse du Soleil dans une vingtaine de kilomètres de diamètre). À l'époque, les bouffées de rayons X, ou sursauts X, sont déjà bien connues : elles résultent d'explosions thermonucléaires à la surface d'étoiles à neutrons qui accumulent du gaz arraché à une étoile compagne. Les théoriciens imaginent alors que les sursauts gamma résultent de conditions d'accrétion particulières, conduisant à une explosion plus vio-

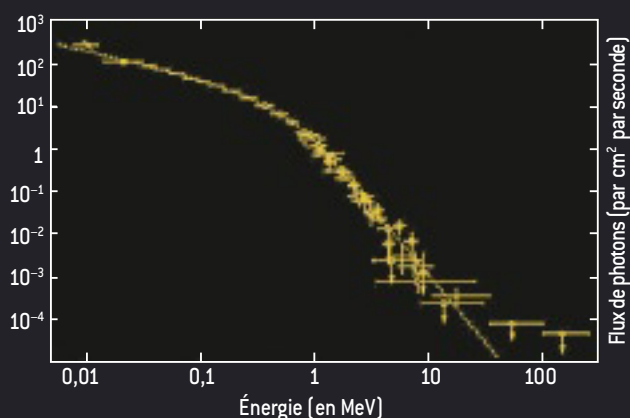
lente. La luminosité déduite de ce modèle implique que les sursauts observés doivent se produire à quelques centaines d'années-lumière tout au plus. À cette échelle, la structure de disque de notre galaxie n'est pas perceptible, ce qui pouvait expliquer la distribution apparemment uniforme (ou « isotrope ») des sursauts sur le ciel. À cet égard, le lancement du satellite CGRO (Observatoire Compton de rayons gamma) en 1991, et notamment du détecteur BATSE (*Burst and transient source experiment*, ou Expérience pour les sources transitoires et les sursauts), soulevait de grands espoirs. Plus sensible que ses prédécesseurs, cet instrument devait détecter des sursauts à plus grande distance et on pensait qu'il révélerait leur origine galactique en montrant une accumulation de sources dans le plan de la Voie lactée.

Après seulement quelques mois d'observation – au rythme d'un sursaut détecté chaque jour en moyenne –, il devint clair que la distribution des sources sur le ciel restait isotrope. Ce résultat signalait l'arrêt de mort de presque tous les modèles théoriques développés jusque-là ! Au terme

1. Le satellite *Swift* a été lancé en novembre 2004 pour étudier les sursauts gamma, de colossales explosions cosmiques. Sa capacité à localiser rapidement et avec précision les sursauts après leur détection a permis de grands progrès dans la compréhension de ces phénomènes.



2. Le profil temporel d'un sursaut gamma présente en général une succession de pics de luminosité enchevêtrés. Parfois, il n'y en a qu'un seul. L'ensemble dure de quelques millisecondes à quelques centaines de secondes. En deçà de deux secondes, il s'agit de sursauts « courts » ; et de sursauts « longs » au-delà. Les deux n'ont pas la même origine.



3. Le spectre d'un sursaut gamma est caractéristique. Le flux de photons émis varie en fonction de leur énergie selon une loi de puissance qui présente une « cassure ». Vers les hautes énergies, la pente s'accroît. L'énergie caractéristique des photons d'un sursaut est de l'ordre de quelques centaines de kiloélectronvolts.

de neuf ans de fonctionnement, BATSE avait détecté 2 704 sursauts répartis de façon uniforme sur la voûte céleste. Les sursauts ne se produisent donc pas dans notre galaxie. D'où proviennent-ils alors ? La localisation des sursauts sur le ciel étant peu précise – quelques degrés, soit plusieurs fois la taille de la Lune –, il était impossible de chercher d'éventuelles contreparties dans le domaine visible aux sources gamma à l'aide de télescopes au sol. Or seules de telles contreparties optiques permettent de mesurer leur distance. Deux théories s'affrontaient ainsi : selon la première, les sursauts se produisaient dans un grand halo sphérique centré sur notre galaxie ; selon la seconde, ils se trouvaient à des distances cosmologiques, au bas mot plusieurs milliards d'années-lumière. La communauté resta divisée jusqu'à la découverte des premières contreparties en 1997.

Lancé un an auparavant, le satellite italo-néerlandais *Beppo-SAX* (satellite pour l'astronomie X *Beppo*) était équipé d'un détecteur gamma moins sensible que BASTE, mais aussi de caméras X d'une précision de quelques minutes d'arc dans la localisation des sources. Le 28 février 1997, après la détection d'un sursaut par le capteur gamma, le satellite fut réorienté pour pointer les caméras X vers cet événement, si bien qu'une position plus précise de la source a pu être transmise au sol huit heures seulement après la détection du sursaut. L'équipe d'Ed van den Heuvel, de l'Université d'Amsterdam, explora alors la région avec le télescope de quatre mètres *William Herschel* situé aux Canaries. Elle découvrit rapidement un objet d'aspect stellaire, jusque-là inconnu et dont l'éclat fut visible pendant près de trois mois en diminuant progressivement. Cette émission, qualifiée de rémanence (*afterglow*, en anglais), apparaissait sur la voûte céleste à proximité d'une galaxie lointaine. Si le sursaut s'était bien produit dans cette galaxie, l'échelle de distance cosmologique s'imposait.

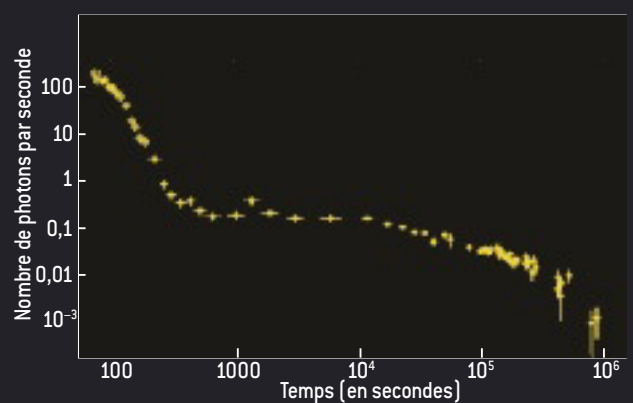
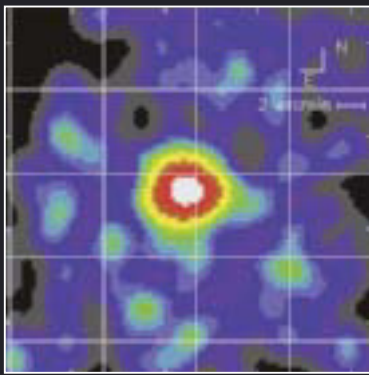
Cependant, un doute subsistait. Il fut définitivement levé par le second sursaut localisé par *Beppo-SAX* en mai de la même année. Sitôt la rémanence découverte, le télescope américain de dix mètres de diamètre *Keck II* parvint à en mesurer le spectre. Celui-ci comportait de fines raies d'ab-

sorption du fer et du magnésium, provenant non du sursaut lui-même, mais plutôt de nuages de gaz absorbant situés le long de la ligne de visée. Une des composantes présentait un décalage vers le rouge dû à l'expansion cosmique égal à $z = 0,835$: le sursaut se trouvait donc au moins à sept milliards d'années-lumière. Depuis lors, le décalage vers le rouge de plus de 100 sursauts gamma a été mesuré. Le plus lointain atteint un décalage de 6,3, ce qui signifie qu'il s'est produit il y a 12,7 milliards d'années. Les sursauts gamma détiennent ainsi le titre de sources les plus brillantes de l'Univers : à une telle distance, seule une luminosité intrinsèque fantastique, équivalente à un milliard de milliards de fois celle du Soleil, peut expliquer les flux observés.

Les sursauts longs sont associés aux supernovae

L'émission rémanente des sursauts gamma a aussi été détectée en ondes radio. De façon générale, le pic d'émission se décale progressivement depuis le domaine X dans les premières heures jusqu'aux ondes radio au bout d'une semaine, en passant par le domaine visible après une journée. En première approximation, l'intensité de la rémanence décroît suivant une loi de puissance dont la pente s'accroît souvent après environ un jour. L'émission n'est cependant pas exempte d'irrégularités, qui se manifestent avec des amplitudes ou des échelles de temps variées. En particulier, certains pics survenant après 30 jours ont été interprétés comme la marque d'une supernova (une explosion d'étoile massive) ayant accompagné le sursaut. Initialement invisible, la lumière de la supernova apparaît après que la rémanence a suffisamment décliné.

L'association entre les sursauts et certains types de supernovae fut confirmée grâce au satellite HETE2, qui prit le relais de *Beppo-SAX* en 2000. Le 29 mars 2003, le satellite détectait un sursaut assez proche ($z = 0,14$, soit 1,7 milliard d'années-lumière environ). Sa rémanence dans le domaine visible, très brillante en raison de sa relative proximité, déclut en



4. Une émission rémanente suit les sursauts gamma. La première a été découverte en rayons X par le satellite *Beppo-Sax* le 28 février 1997, huit heures après la détection du sursaut GRB 970228 (à gauche). Le même jour, une contrepartie optique était découverte par le télescope *Herschel*, aux Canaries (au milieu, flèche). Visible en

rayons X dans les premières heures, l'émission rémanente se décale jusqu'aux ondes radio après plusieurs jours, en passant par le domaine visible. Dans le domaine X, l'émission décroît d'abord fortement, puis connaît un plateau après quelques centaines de secondes, avant de décliner de nouveau rapidement au bout de quelques heures (à droite).

quelques jours en laissant apparaître la forme caractéristique du spectre d'une supernova de type Ic. Une telle supernova correspond à l'explosion d'une étoile de Wolf-Rayet, une étoile massive en fin de vie qui a soufflé ses couches d'hydrogène et d'hélium au point qu'il ne subsiste que son cœur de carbone et d'oxygène au moment de l'explosion. La signature spectroscopique observée dans la rémanence du sursaut GRB 030329 est la preuve indiscutable que l'explosion d'une étoile massive avait accompagné le sursaut, sans qu'il soit possible de préciser si elle le précédait ou le suivait. L'association directe d'un sursaut long avec une supernova a été confirmée dans plusieurs cas, et on pense que c'est la norme.

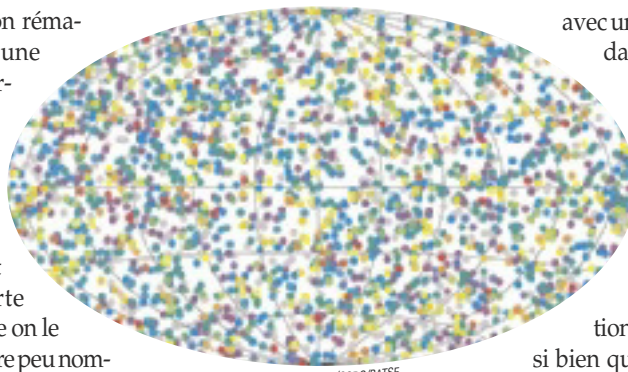
L'observation de l'émission rémanente ne concernait, jusqu'à une époque très récente, que les sursauts longs. Pour des raisons techniques, les caméras X de *Beppo-SAX* et de *HETE 2* échouaient à localiser les sursauts courts. La première rémanence d'un sursaut court ne fut finalement découverte qu'en 2005 par *HETE 2*. Comme on le verra, toutes les suivantes, encore peu nombreuses, ont été détectées par son successeur *Swift*. Leurs caractéristiques suggèrent une origine différente de celle des sursauts longs, probablement la coalescence de deux objets compacts – deux étoiles à neutrons ou une étoile à neutrons et un trou noir. Mais avant de nous pencher sur les sursauts courts, voyons quel est le mécanisme général de production des sursauts gamma.

Malgré les progrès accomplis depuis l'élaboration des premiers modèles, de nombreuses zones d'ombres subsistent. Le scénario décrivant l'émission gamma aujourd'hui le plus discuté est celui dit des « chocs internes ». La variabilité temporelle des sursauts, de l'ordre de la milliseconde, implique une source très petite, telle une étoile à neutrons ou trou noir de quelques masses solaires.

Le flux lumineux gigantesque trouve donc *a priori* naissance dans une région de quelques dizaines de kilomètres d'extension tout au plus. La densité d'énergie y est telle que la matière produit des photons ultra-énergétiques, notamment dans le domaine gamma.

Un problème se pose cependant : une partie des photons gamma devraient s'annihiler en paires d'électrons et de positons. Ces particules, interagissant en retour avec les photons, rendraient la source opaque à son propre rayonnement. Ce paradoxe, dit de compacité, peut être levé si l'on imagine que le gaz qui rayonne se déplace vers l'observateur à une vitesse proche de celle de la lumière. Un photon observé avec une énergie E aura alors été émis avec une énergie propre E' plus petite dans le repère lié au gaz en mouvement. Pour une vitesse suffisamment grande – au moins 0,99995 fois celle de la lumière – E' est 100 fois plus petite que E et très peu de photons atteignent alors l'énergie de 511 kiloélectronvolts nécessaire à la création d'une paire électron-positon, si bien que le gaz reste transparent au rayonnement.

Le sursaut gamma résulte donc d'un flot de matière accéléré jusqu'à des vitesses très proches de celle de la lumière. Quelle est la nature du « moteur central » capable de communiquer de telles vitesses ? Les mécanismes d'accélération restent assez mal compris. Il s'agit peut-être de processus thermiques (expansion d'un gaz chauffé à très haute température) ou magnétiques (plasma guidé par un champ magnétique ancré dans la source). Le moteur



NASA/CGRO/BATSE

5. La distribution des sursauts gamma sur la voûte céleste est isotrope, comme l'ont révélé les observations de l'instrument BATSE du satellite CGRO dans les années 1990 (ci-dessus). Ce résultat exclut une origine galactique. On sait aujourd'hui que les sursauts gamma se produisent à des distances cosmologiques.

Anatomie d'un sursaut gamma long

Une étoile de Wolf-Rayet, un astre massif en fin de vie, s'effondre sur elle-même.

L'enveloppe de gaz est éjectée lors de l'explosion de l'étoile en supernova.

Jet

Disque d'accrétion

Rayonnement gamma

Trou noir

Le cœur s'effondre en un trou noir. Le gaz restant tombe dessus via un disque d'accrétion. Une partie de cette matière est éjectée à des vitesses proches de celle de la lumière, engendrant le rayonnement gamma observé.

central pourrait ainsi être un trou noir de quelques masses solaires entouré d'un disque d'accrétion (d'autres objets ont aussi été envisagés, telles des étoiles à neutrons ultra-magnétisées ou métastables). Dans le cas des sursauts longs, cette configuration résulterait de l'effondrement du cœur d'une étoile massive en rotation rapide, les régions périphériques donnant naissance à un disque, car leur moment angulaire élevé les empêche de tomber directement dans le trou noir. Dans les sursauts courts, le trou noir et son disque se formeraient à la suite de la coalescence des deux objets compacts, par exemple deux étoiles à neutrons ou une étoile à neutrons et un trou noir. Dans les deux cas, la source d'énergie est soit l'accrétion de la matière du disque par le trou noir, ou l'extraction directe de l'énergie de rotation du trou noir par le champ magnétique.

Dans le modèle standard des sursauts gamma, il n'est cependant pas nécessaire de comprendre en détail la physique très complexe de la source pour étudier l'origine de l'émission gamma proprement dite. Supposons simplement qu'un jet de matière est, d'une façon ou d'une autre, éjecté à une vitesse proche de celle de la lumière. La vitesse de l'écoulement n'a *a priori* aucune raison d'être parfaitement uniforme tout au long du jet. Diverses instabilités opérant sur des échelles de temps variées tendent en effet à rendre la distribution des vitesses inhomogène. Les couches les plus rapides vont alors rattraper les couches plus lentes. Si par exemple une bouffée de gaz se déplaçant à 99,9994 pour cent de la vitesse de la lumière est émise une seconde après une autre filant à 99,995 pour cent de la vitesse de la lumière, elle la rattrapera à environ six milliards de kilomètres de la source. Dans le référentiel de la première couche, la seconde la percute à 80 pour cent de la vitesse de la lumière.

Les sursauts prennent naissance au sein de jets de matière

Une partie de l'énergie dissipée dans ce choc est transférée à des électrons du milieu ambiant, qui rayonnent en se déplaçant à grande vitesse dans le champ magnétique local, amplifié dans la collision. C'est ce rayonnement, dit synchrotron, qui constitue l'émission gamma observée. En pratique, l'éjection de matière est continue, et pour en reproduire fidèlement l'évolution, il faut considérer un grand nombre de couches. De multiples chocs se produisent et s'entremêlent, et la résultante de toutes leurs contributions conduit au profil observé du sursaut. Le modèle des chocs internes, introduit en 1994 par Martin Rees, de l'Université de Cambridge, et Peter Meszaros, à l'Université Pennstate, en Pennsylvanie, et développé par Frédéric Daigne et moi-même à l'Institut d'astrophysique de Paris à la fin des années 1990, montre un accord satisfaisant avec les données observationnelles sur l'émission gamma. Cependant l'efficacité globale du processus est faible, inférieure à dix pour cent. En d'autres termes, il faut que soit injectée dans le jet relativiste une énergie plus de dix fois supérieure à celle, déjà très élevée, que l'on observe dans le domaine gamma. Cela impose une forte contrainte énergétique sur la source centrale.

Dider Florentz

Anatomie d'un sursaut gamma court

Comment se produit ensuite l'émission rémanente ? Jusqu'à récemment, cela semblait être la partie la mieux comprise du modèle. On l'expliquait par le freinage de la matière éjectée au cours de sa progression dans le milieu environnant. Cette idée avait été proposée dès 1976 dans un autre contexte par Roger Blandford et Christopher McKee, alors à l'Université de Berkeley, en Californie. Elle généralisait à des vitesses relativistes le mécanisme d'évolution des restes de supernovae proposé par Leonid Sedov, de l'Université de Moscou, en 1959. La matière éjectée, agissant comme un piston, provoque la formation d'une onde de choc qui balaie le milieu extérieur en avant du jet – on parle de choc avant – alors que simultanément une onde de choc « retour » traverse le jet en sens inverse en le freinant. Derrière le choc avant, les différentes grandeurs physiques (vitesse, densité, température ou facteur de contraction relativiste) varient avec le temps selon une loi de puissance. Cette solution reproduit assez bien les données d'observation, en supposant que c'est le rayonnement synchrotron des électrons accélérés dans le choc avant qui produit le flux rémanent observé.

Les surprises de Swift

Le lancement du satellite *Swift*, en novembre 2004, allait forcer les astrophysiciens à revoir ce modèle. Plus sensible que *Beppo-SAX* ou *HETE 2*, *Swift* est en outre capable de localiser les sursauts en temps réel et de se réorienter en à peine quelques minutes pour braquer dans leur direction un télescope à rayons X d'une grande sensibilité (XRT) et un télescope ultraviolet et visible (UVOT). Autant d'atouts pour localiser rapidement un sursaut et observer la rémanence associée dès les premières minutes, qui ont permis d'accomplir des progrès importants, mais ont aussi apporté leur lot de surprises.

Avant le lancement de *Swift*, l'évolution de la rémanence dans les premières heures suivant la fin du sursaut était très mal connue. On imaginait pouvoir simplement extrapoler « en ligne directe » les profils d'émission pour raccorder la rémanence au sursaut. Dès les premières observations de *Swift*, il est apparu que le phénomène est bien plus complexe. L'évolution d'une rémanence typique connaît en effet plusieurs phases. L'émission commence d'abord par décroître fortement pendant quelques centaines de secondes, puis elle connaît un plateau durant lequel la décroissance est faible, avant de redevenir plus prononcée au bout de quelques heures. Avant les observations de *Swift*, seule cette dernière phase était connue. Pour compliquer encore davantage la situation, des pics parfois intenses se superposent à ce comportement général.

Le modèle classique ne suffit plus à expliquer cette évolution. Si la rapide décroissance initiale peut être interprétée comme la fin de la contribution des chocs internes, le plateau intermédiaire est plus problématique. La lenteur du déclin durant cette phase a été expliquée par une activité tardive de la source, celle-ci injectant en continu de l'énergie dans l'onde de choc avant. De la même manière, une activité prolongée de la source pourrait conduire à de nouveaux chocs internes, à l'origine des pics sporadiques.

Toutefois, cette interprétation soulève plusieurs questions. Quelle source pourrait rester active plusieurs heures après l'événement catastrophique initial ? Par ailleurs, si

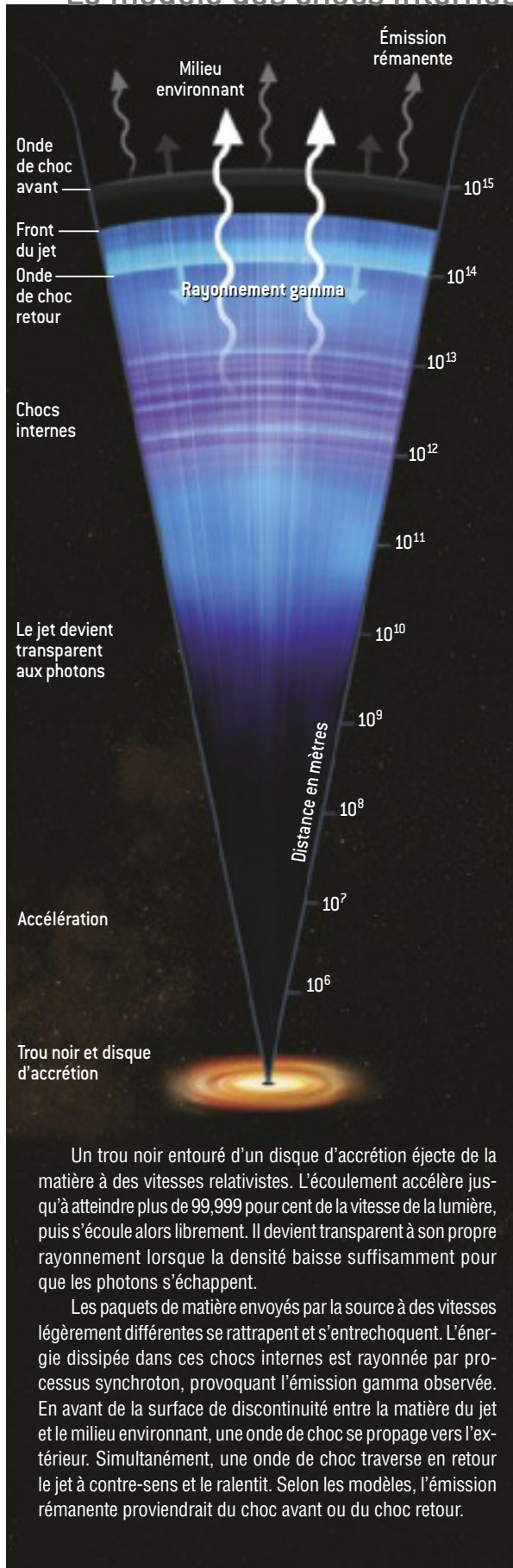
Deux objets compacts, par exemple deux étoiles à neutrons ou une étoile à neutrons et un trou noir, gravitent l'un autour de l'autre.

Au fil du temps, les deux astres perdent de l'énergie gravitationnelle et se rapprochent.

Les deux objets finissent par tomber l'un sur l'autre. Ils fusionnent alors en une fraction de seconde, donnant naissance à un trou noir entouré d'un disque d'accrétion.

Comme pour les sursauts longs, le rayonnement gamma provient d'un jet issu du disque. Cependant, le disque ne représente que quelques dixièmes de masse solaire contre quelques dizaines pour les sursauts longs. Le phénomène est donc beaucoup plus bref.

Le modèle des chocs internes



une grande partie de l'énergie de l'événement est injectée après que les chocs internes ont cessé, comment l'émission gamma observée est-elle engendrée à partir d'une énergie initiale moindre ? D'autant que, comme nous l'avons évoqué, l'efficacité des chocs internes est assez faible.

Ces difficultés ont conduit à des remises en question du modèle standard. Récemment, Frédéric Daigne, moi-même et notre étudiant en thèse Franck Genet, d'une part, et Andrei Beloborodov et son étudiant Lucas Uhm, de l'Université Columbia à New York, d'autre part, avons proposé simultanément que la rémanence résulte non pas du choc avant, mais plutôt du choc retour qui traverse l'éjectat au cours de sa décélération. Cette hypothèse semble pouvoir expliquer le plateau observé dans la rémanence précoce et peut-être même aussi la présence des pics sans faire intervenir une activité tardive de la source. Il reste cependant encore à la tester.

Les sursauts courts résultent de la fusion de deux astres compacts

L'autre avancée importante permise par *Swift* concerne les sursauts courts. Du fait de leur brièveté, ceux-ci sont plus difficiles à localiser que les sursauts longs, et *Beppo-SAX* n'avait pu déterminer la position d'aucun sursaut avec précision, ce qui interdisait toute recherche de l'émission rémanente associée. Or, comme nous l'avons rappelé, c'est par l'étude de la rémanence que peuvent être déduites la distance du sursaut et la nature de l'objet responsable de l'explosion. La première rémanence d'un sursaut court, observée en 2005 par HETE 2, était située au voisinage d'une galaxie elliptique. Si le sursaut s'est bien produit dans cette galaxie, l'objet parent ne peut pas être une étoile massive comme c'est le cas pour les sursauts longs. Les étoiles massives vivent en effet moins de dix millions d'années et ne se trouvent que dans des galaxies où la formation stellaire est active. Or les galaxies elliptiques sont de vieilles galaxies assoupies, peuplées d'étoiles âgées. *Swift* a par la suite réussi à observer une dizaine de rémanences de sursauts courts, qui ont confirmé la découverte. En particulier, aucune supernova sous-jacente n'a été découverte accompagnant ces sursauts courts, que ce soit par spectroscopie ou sous forme de pics tardifs dans la rémanence.

Quelle est alors la source des sursauts gamma courts ? L'ensemble de ces résultats exclut l'hypothèse de la mort d'une étoile massive et suggère plutôt la coalescence de deux objets compacts. Les systèmes binaires serrés constitués de deux étoiles à neutrons ou d'une étoile à neutrons et d'un trou noir qui gravitent l'un autour de l'autre, évoluent en effet lentement jusqu'à la fusion. Les deux astres perdent progressivement de l'énergie liée au mouvement orbital par l'émission d'ondes gravitationnelles, et de ce fait, se rapprochent inexorablement. Lorsqu'ils sont suffisamment proches, ils fusionnent en quelques fractions de seconde pour donner naissance à un trou noir entouré d'un disque d'accrétion épais, qui produit un sursaut gamma sans qu'il y ait de supernova associée. Cette idée a été proposée en 1992 par Ramesh Narayan et Tsvi Piran, de l'Université de Cambridge, et Bohdan Paczynski, de l'Université de Princeton. La quantité de matière disponible est ici de quelques dixièmes de masse solaire contre quelques dizaines dans le cas des sur-

Didier Florentz

sauts longs, ce qui explique la brièveté de ce type de sursaut. Le processus de rapprochement peut pour sa part durer plusieurs milliards d'années, ce qui introduit un délai très long entre la formation du système binaire et l'événement catastrophique final. Cela expliquerait que les sursauts courts peuvent se produire dans les galaxies elliptiques.

Enfin, *Swift*, a dévoilé des sursauts gamma de plus en plus lointains. La rapidité du satellite à localiser avec précision les sursauts a favorisé la recherche des contreparties optiques par les grands télescopes terrestres. La distance de plus de 80 sursauts a ainsi été mesurée en moins de trois ans. Avant le lancement de *Swift*, la valeur moyenne des décalages était de 1,2 avec un maximum à 4,5. Avec *Swift*, la moyenne est passée à 2,6 et le maximum a été porté à 6,3. Cette valeur est presque comparable à celle des quasars les plus distants (des galaxies abritant en leur centre un trou noir supermassif). Le record de l'objet le plus lointain sera peut-être bientôt détenu par un sursaut gamma !

Les sursauts gamma sont ainsi devenus des outils de choix pour sonder l'Univers profond. Comme pour les quasars, l'analyse spectroscopique des rémanences permet de déduire la nature et la distribution des objets présents le long de la ligne de visée, qui produisent des raies d'absorption décalées selon leur éloignement. En particulier, cette méthode permet d'étudier le milieu interstellaire et les propriétés de la galaxie hôte du sursaut jusqu'à de grandes distances alors que la galaxie est elle-même trop faible pour être directement observée. Cependant le déclin rapide de la rémanence impose d'obtenir un spectre dans des délais très courts. L'Observatoire européen austral (ESO) développe dans ce but un

spectrographe rapide appelé *X-shooter* qui sera installé l'an prochain sur le très grand télescope européen VLT.

Au cours des 15 dernières années, notre conception des sursauts gamma a profondément changé. Cet élan va se poursuivre en 2008 avec le lancement du satellite GLAST, qui observera les sursauts à très haute énergie (de la dizaine de mégaélectronvolts au gigaélectronvolt). L'expérience EGRET, sur le satellite CGRO, a détecté des sursauts et leurs rémanences dans ce domaine d'énergie dès les années 1990, mais sa sensibilité limitée n'a pas permis d'étude détaillée. Beaucoup plus performant, GLAST ouvrira cette nouvelle fenêtre spectrale et posera de nouvelles contraintes sur les processus d'émission. À l'horizon 2012 sera lancé le satellite franco-chinois SVOM/ECLAIRS. Moins sensible que *Swift*, SVOM couvrira en revanche un domaine spectral plus étendu – de quelques kiloélectronvolts à quelques mégaélectronvolts – et sera optimisé pour la recherche de sursauts lointains. Entre le défi que pose la physique de ces objets extrêmes et leur intérêt pour étudier les confins de l'Univers, l'attrait exercé par les sursauts gamma n'est pas près de faiblir.

Robert MOCHKOVITCH est chercheur à l'Institut d'astrophysique de Paris.

N. GEHRELS, J. K. CANNIZZO, et J.P. NORRIS, *Gamma-ray bursts in the Swift era*, in *New Journal of Physics*, vol. 9, n° 2, pp. 37, 2007.

E. NAKAR, *Short-hard gamma-ray bursts*, in *Physics Reports*, vol. 442, pp. 166-236, 2007, ou astro-ph/0701748.

T. PIRAN, *The physics of gamma-ray bursts*, in *Reviews of Modern Physics*, vol. 76, pp. 1143-1210, 2004, ou astro-ph/0405503.

Auteur & Bibliographie

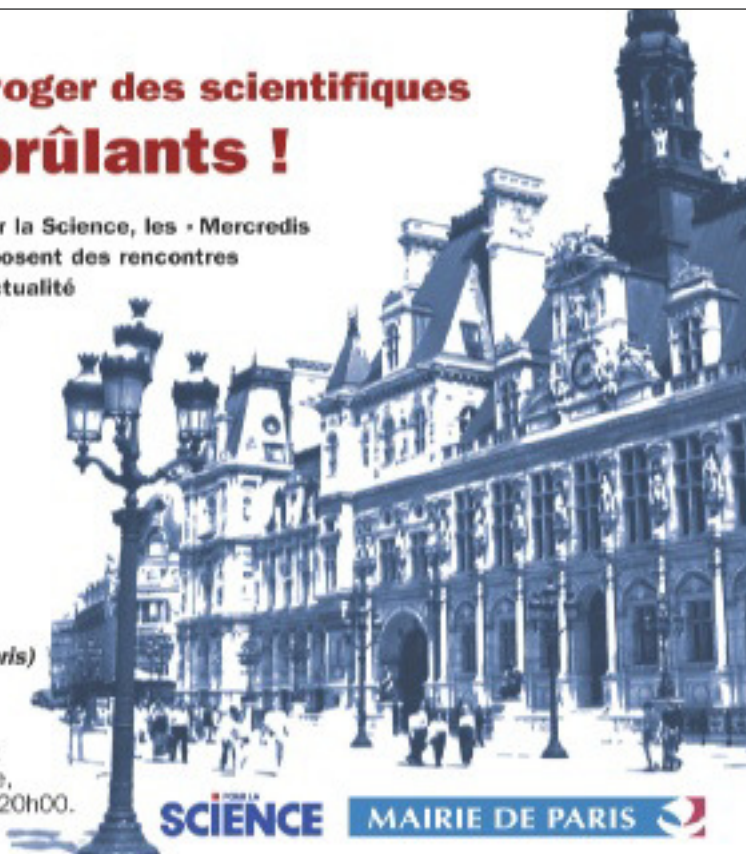
Venez interroger des scientifiques sur des sujets brûlants !

Organisés par la Mairie de Paris et Pour la Science, les « Mercredis Pour la Science à l'Hôtel de Ville » proposent des rencontres avec des spécialistes qui éclairent l'actualité scientifique de façon complémentaire.

13 février Les extinctions massives d'espèces

Avec : **Éric Buffetaut** (Paléontologue au CNRS)
et **Vincent Courtillot** (Directeur de l'Institut de physique du Globe de Paris)

Ouvertes à tous, ces rencontres auront lieu dans l'auditorium de l'Hôtel de Ville, 3, rue Lobau, 75004 Paris, de 18h30 à 20h00. Entrée libre.



POUR LA SCIENCE

MAIRIE DE PARIS

Éviter la mort des neurones

Dans la maladie de Charcot ou sclérose latérale amyotrophique, les neurones moteurs qui innervent les muscles dégénèrent progressivement : les patients finissent par être paralysés. Les neurobiologistes précisent les mécanismes de cette maladie et plusieurs traitements sont à l'étude.

Patrick Aebischer • Ann Kato

La maladie de Charcot est aussi nommée sclérose latérale amyotrophique (SLA), et, aux États-Unis, maladie de Lou Gehrig. En effet, on a diagnostiqué cette maladie chez ce joueur de baseball des *New York Yankees* en 1939 ; il en est mort deux ans plus tard. C'est une maladie neuromusculaire attaquant progressivement les neurones dits moteurs (ou motoneurones) qui relient le cerveau et la moelle épinière à l'ensemble des muscles du corps. Quand ces motoneurones meurent, le cerveau ne peut plus contrôler les muscles ; aux stades avancés de la maladie, les patients sont paralysés.

La sclérose latérale amyotrophique a été décrite dès 1869 par Jean-Martin Charcot qui fonda l'école de neurologie de l'Hôpital de la Salpêtrière, à Paris. Pourtant, près de 140 ans après, on ignore encore pourquoi les motoneurones dégénèrent. Elle fut longtemps considérée comme une maladie rare, mais en fait, elle touche 8 000 personnes en France, 6 000 en Allemagne, 5 000 en Grande-Bretagne et 30 000 aux États-Unis. La sclérose latérale amyotrophique apparaît en général entre 40 et 70 ans. Plusieurs personnalités ont été touchées par la maladie : l'acteur britannique David Niven, le compositeur russe Dmitri Chostakovitch et le leader chinois Mao Tsé-toung par exemple. Un nombre important de patients sont des joueurs de football italiens, des vétérans de la guerre du Golfe et des résidents de l'île de Guam (de l'archipel des Mariannes, en Micronésie), mais on ignore pourquoi.

Les patients décèdent souvent trois à cinq ans après le diagnostic... à l'exception de Stephen Hawking, le physicien de l'Université de Cambridge, qui vit avec une sclérose latérale amyotrophique depuis plus de 40 ans et qui continue à faire des recherches en cosmologie et en physique quantique malgré son handicap (voir la figure 3). Comment et pourquoi les motoneurones dégénèrent-ils ? Depuis quelques années, les neurobiologistes ont fait des progrès considérables et comprennent mieux les mécanismes de dégénérescence de ces

neurones. Ils espèrent ainsi développer des traitements susceptibles de retarder l'évolution de la maladie, voire d'en empêcher le déclenchement. Nous aborderons ici quelques-unes des thérapies envisagées.

Une paralysie progressive de presque tous les muscles

Que signifie sclérose latérale amyotrophique ? *Amyotrophique* vient du grec *a* privatif, *myo* muscle et *trophique* nourriture. Ainsi, les muscles des patients ne sont plus nourris, de sorte qu'ils s'atrophient. *Latérale* renvoie à l'endroit de la moelle épinière où sont situées les parties des neurones qui meurent : la racine latérale. Quand cette région dégénère, elle durcit ; *sclérose* signifie qui durcit. L'aspect le plus terrible de la maladie est que les fonctions cérébrales supérieures du patient restent intactes ; le patient est parfaitement conscient de son état, et il assiste impuissant à l'agonie de son corps.

La forme la plus répandue de la maladie de Charcot est dite sporadique, c'est-à-dire qu'elle frappe au hasard, n'importe qui, n'importe où ; on en ignore les causes. La sclérose latérale amyotrophique familiale est une forme particulière qui se transmet de génération en génération, mais seuls cinq à dix pour cent des malades sont concernés.



moteurs



1. Dans la maladie de Charcot, les motoneurones commencent à dégénérer aux extrémités de l'axone (en bleu). Les anomalies se propagent ensuite vers le corps cellulaire (en orange en haut à droite). C'est un phénomène de mort dit rétrograde.

Quels sont les premiers symptômes de la maladie ? Ils varient d'une personne à l'autre ; en général, le patient laisse tomber des objets et a l'impression que ses bras et ses jambes fatiguent vite ; il souffre de crampes musculaires et de petites contractions que l'on nomme secousses musculaires. Il a des difficultés pour parler, marcher et utiliser ses mains dans des activités quotidiennes, tels la toilette ou l'habillage. À mesure que l'affaiblissement et la paralysie se propagent aux muscles du tronc, il n'arrive plus à déglutir, mâcher... puis respirer. Quand les muscles de la respiration sont atteints, le patient ne survit qu'avec une ventilation assistée.

Étant donné que la sclérose latérale amyotrophique ne détruit que les neurones moteurs, les cinq sens – la vision, le toucher, l'audition, le goût et l'odorat – ne sont pas altérés. Qui plus est, pour des raisons inconnues, les motoneurones responsables des mouvements des yeux et de la vessie sont épargnés pendant longtemps. Par exemple, S. Hawking contrôle toujours ses muscles oculaires ; d'ailleurs, à une époque, il communiquait en élevant un sourcil quand son assistant désignait des lettres sur un écran. Il bouge aussi deux doigts de sa main droite, et utilise maintenant un synthétiseur de parole activé par un interrupteur manuel.

Aujourd'hui, le seul traitement contre la sclérose latérale amyotrophique est le riluzole : cette molécule prolonge la vie de plusieurs mois, car elle ralentit la libération d'agents chimiques délétères qui endommagent les neurones moteurs.

Mais quelles sont les causes de la sclérose latérale amyotrophique ? Les chercheurs ont proposé une pléthore de responsables : des agents infectieux, un système immunitaire défectueux, des gènes, des substances toxiques, un déséquilibre chimique dans le corps et une alimentation de mauvaise qualité. Bien que l'on ignore encore ce qui déclenche la maladie chez la plupart des personnes atteintes, on a fait un grand pas en avant en 1993 : un consortium de généticiens et de cliniciens a découvert un gène responsable de l'une des formes héréditaires de la sclérose latérale amyotrophique, qui représente environ deux pour cent des cas. Le gène incriminé code une enzyme nommée superoxyde dismutase (SOD1) qui protège les cellules des dommages provoqués par les radicaux libres – des molécules très réactives, et par conséquent nocives, produites lors de processus métaboliques normaux.

On a ensuite identifié plus de 100 mutations différentes du gène codant SOD1, ces mutations engendrant une sclérose latérale amyotrophique. Mais on ignore toujours comment des altérations de cette enzyme, présente dans toutes les cellules de l'organisme, peuvent produire des lésions ne touchant qu'un seul type de neurones. Au départ, on pensait que la toxicité était due au fait que les neurones perdaient leur capacité à se défendre contre les radicaux libres. Mais aujourd'hui, plusieurs données suggèrent que l'enzyme SOD1 mutée acquiert une propriété destructrice, un phénomène nommé « gain de fonction ».

Grâce à cette découverte, les neurobiologistes ont inséré des ADN modifiés dans le génome de souris

de laboratoire afin de créer des lignées de rongeurs qui présentent une maladie des neurones moteurs semblable à la sclérose latérale amyotrophique. Ainsi, comme on a pu suivre l'évolution de la maladie de ces rongeurs, le nombre de recherches et de publications sur la maladie de Charcot a explosé. Puis on a trouvé d'autres gènes qui entraînent la mort des motoneurones et on a développé d'autres modèles animaux. Finalement, on commence à comprendre le mécanisme de dégénérescence de ces neurones dans la sclérose latérale amyotrophique.

Une particularité du motoneurone est que son axone est très long : cette branche principale qui part du corps cellulaire peut faire plus de un mètre chez un adulte (voir l'encadré ci-dessous). À son extrémité, l'axone se divise en petites ramifications qui se projettent à la surface du muscle comme les dents d'un râteau. Chaque connexion entre le nerf et le muscle est appelée une synapse. En libérant de petites molécules nommées neurotransmetteurs, le nerf provoque la contraction du muscle.

On a longtemps cru que le neurone moteur et ses diverses parties mouraient simultanément. Mais on sait aujourd'hui que les différents compartiments du motoneurone dégèrent selon divers mécanismes. Le corps cellulaire, qui contient le noyau du neurone, meurt souvent par un mécanisme d'apoptose : il s'autodétruit en se fragmentant et en empaquetant les débris dans de petits sacs facilement éliminés. L'axone meurt selon un autre mécanisme, et la synapse, vraisemblablement, par un troisième. Si le neurone commence à mourir à partir de la région terminale de l'axone, il s'agit d'une mort rétrograde. S'il meurt à partir du corps cellulaire, c'est une mort antérograde. De plus en plus

de données suggèrent que dans le cas de la sclérose latérale amyotrophique, les premiers signes de dégénérescence des motoneurones apparaissent au niveau de la synapse et de l'axone ; la mort est donc rétrograde (voir la figure 1). En conséquence, on cherche avant tout des molécules capables de protéger les axones.

Comment les neurones moteurs meurent-ils ?

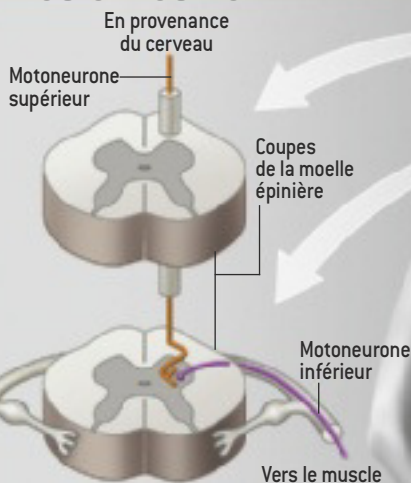
Nous l'avons déjà mentionné : certains neurones moteurs sont plus vulnérables que d'autres. On commence à comprendre pourquoi. Récemment, Pico Caroni et ses collègues, de l'Institut Friedrich Miescher à Bâle, en Suisse, ont abordé ce problème en cherchant comment les impulsions nerveuses se propagent des motoneurones aux muscles d'une souris. Ils ont créé une souris génétiquement modifiée qui exprime des marqueurs fluorescents dans les axones de certains motoneurones. Ainsi, les neurones se connectant aux muscles peuvent être regroupés en trois catégories. Les premiers innervent les fibres musculaires des bras et des jambes, dites fatigables, à contraction rapide qui permettent les mouvements rapides, et consomment beaucoup d'énergie. On les nomme FF (pour *fast-fatigable*). La deuxième classe de motoneurones contrôle les fibres musculaires à contraction rapide, mais qui résistent à la fatigue (FR pour *fast-resistant*) ; elles sont nombreuses dans les petits muscles des membres. Et le troisième groupe innerve les fibres à contraction lente (S pour *slow*) qui sont sollicitées en permanence, par exemple les muscles du tronc impliqués dans le maintien de la posture.

Quand l'équipe de P. Caroni a examiné le marqueur fluorescent chez les souris SOD1 génétiquement vulnérables à la sclérose latérale amyotrophique, elle a découvert que les neurones contrôlant les fibres FF dégèrent en premier, tandis que les neurones FR sont touchés plus tard, et les neurones S restent intacts jusqu'aux stades avancés de la maladie. Cette évolution ressemble beaucoup à la progression des symptômes chez les patients atteints de sclérose latérale amyotrophique, ce qui renforce la similitude entre les formes animale et humaine de la maladie.

Joshua Sanes et Jeff Lichtman, de l'Université Harvard, ont aussi étudié les axones fluorescents des motoneurones chez des souris atteintes de sclérose latérale amyotrophique. Ainsi, ils ont distingué les motoneurones qui dégèrent de leurs voisins qui tentent de se régénérer. Il existe donc deux types de neurones dans les mêmes circuits moteurs des souris : les « perdants », qui se fragmentent tellement que leur liaison au muscle est interrompue, et les « compensateurs », qui émettent de nouveaux axones vers les muscles (pour remplacer ceux perdus). Reste à trouver les facteurs qui maintiennent les neurones compensateurs en vie, pour peut-être développer de nouveaux traitements.

En outre, on a montré qu'il est nécessaire de protéger non seulement le corps cellulaire, mais aussi

Les cibles de la maladie de Charcot



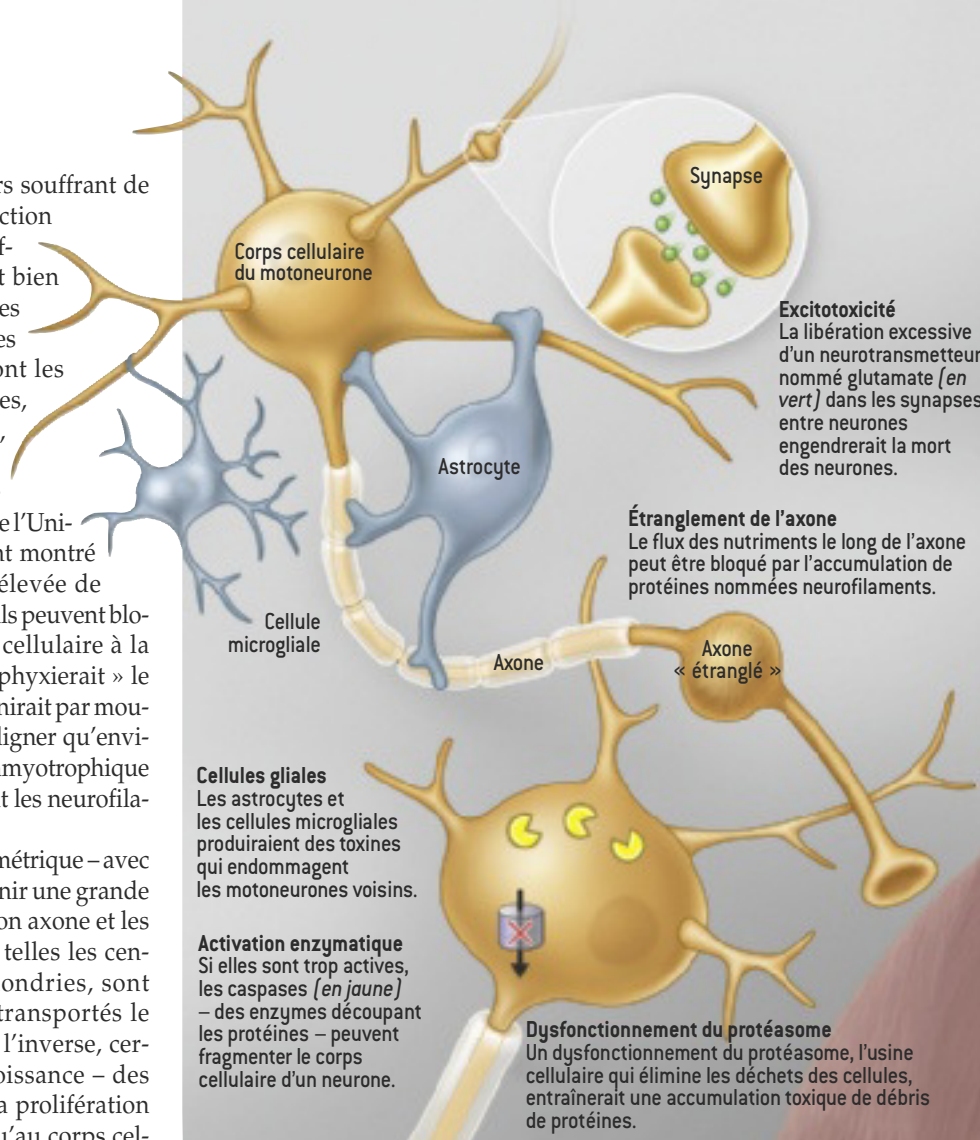
L'évolution de la maladie de Charcot varie beaucoup d'un patient à l'autre. Toutefois, la maladie touche à la fois les motoneurones supérieurs et inférieurs. Les corps cellulaires des motoneurones supérieurs se trouvent dans le cortex moteur du cerveau [1], et leurs axones atteignent soit le tronc cérébral [2], soit la moelle épinière [3]. Ils se connectent alors aux motoneurones inférieurs [4, en violet] dont les axones transmettent le signal aux muscles [5]. C'est ainsi que les muscles sont contrôlés.

l'axone, pour prolonger la vie des rongeurs souffrant de sclérose latérale amyotrophique ; la protection du seul corps cellulaire est insuffisante. Différents mécanismes moléculaires seraient bien impliqués dans la mort des différentes parties du neurone moteur. Par exemple, les protéines qui donnent sa forme rigide à l'axone sont les neurofilaments. Pour des raisons inconnues, plus un axone contient de neurofilaments, plus son diamètre est grand et plus il est vulnérable à la dégénérescence dans la sclérose latérale amyotrophique. Jean-Pierre Julien, de l'Université Laval à Québec, et ses collègues ont montré que lorsqu'une quantité anormalement élevée de neurofilaments s'accumule dans les axones, ils peuvent bloquer le transfert de nutriments du corps cellulaire à la synapse. Cette anomalie dans l'axone « asphyxierait » le corps cellulaire des neurones moteurs, qui finirait par mourir (voir la figure 2). Il est intéressant de souligner qu'environ un pour cent des cas de sclérose latérale amyotrophique est dû à des mutations des gènes qui codent les neurofilaments.

Étant donné que le motoneurone est asymétrique – avec un long axone –, le corps cellulaire doit fournir une grande quantité d'énergie pour maintenir en vie son axone et les synapses. Plusieurs organites cellulaires, telles les centrales énergétiques que sont les mitochondries, sont fabriqués dans le corps cellulaire, puis transportés le long de l'axone jusqu'à la terminaison ; à l'inverse, certaines substances, tels les facteurs de croissance – des protéines qui stimulent la maturation et la prolifération cellulaires –, transitent de la synapse jusqu'au corps cellulaire. Tout ce trafic intracellulaire nécessite des moteurs moléculaires qui sont liés comme des roues dentées à de minuscules rails composés de microtubules. La rupture de l'un de ces éléments peut perturber le fonctionnement du neurone moteur. Par exemple, plusieurs équipes ont montré que des mutations des gènes codant les moteurs moléculaires peuvent tuer les motoneurones.

De surcroît, on pensait que les neurones moteurs des patients souffrant de sclérose latérale amyotrophique provoquaient leur propre destruction... Ce ne serait peut-être pas le cas. On a récemment montré que les cellules gliales avoisinantes, qui soutiennent les neurones et qui les nourrissent, interviendraient dans la genèse des lésions. Ainsi, chez les souris, la sclérose latérale amyotrophique ne se développe pas si l'enzyme SOD1 est mutée seulement dans les neurones ou dans les cellules gliales. De plus, l'équipe de Don Cleveland, à l'Université de Californie, à San Diego, a montré que des cellules gliales saines peuvent protéger des motoneurones malades, et, inversement, que des cellules gliales malades déclenchent la dégénérescence de neurones moteurs sains. Par conséquent, les neurones moteurs et les cellules gliales voisines s'associent pour engendrer la maladie de Charcot.

Voilà quelques progrès récents dans la compréhension de la maladie ; pourra-t-on désormais mettre au point des traitements efficaces ? On connaît déjà certaines molécules qui peuvent protéger les axones des motoneurones. L'une d'elles est une protéine nommée facteur neurotrophique ciliaire (ou CNTF) qui est essentielle à la



2. Comment la sclérose latérale amyotrophique tue ? On connaît aujourd'hui plusieurs mécanismes qui seraient responsables de la dégénérescence des motoneurones chez les personnes atteintes de la maladie de Charcot. Différentes parties du neurone – le corps cellulaire, l'axone, etc. – meurent selon des processus différents. Les cellules gliales voisines, qui normalement protègent les neurones, seraient aussi impliquées dans le déclenchement de la maladie.

survie des motoneurones et des neurones sensoriels. D'autres molécules, tel le facteur neurotrophique dérivé des cellules gliales (ou GDNF), protègent le corps cellulaire des neurones de l'apoptose, mais n'ont aucun effet sur leur axone...

Vivre plus longtemps

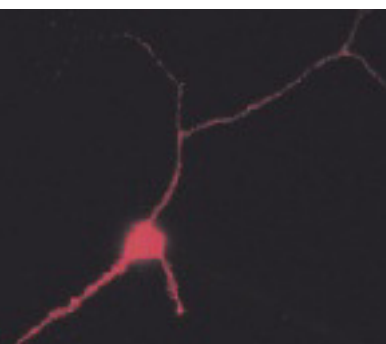
L'équipe de Jeffrey Milbrandt, de l'Université Washington, à Saint-Louis, a travaillé avec une souris mutante nommée Wld^S chez laquelle il existe un mécanisme naturel de protection des axones, car, dans l'ADN du rongeur, deux gènes normalement distincts sont fusionnés. Cette séquence génétique code une protéine chimérique composée de deux parties : un peptide – un fragment protéique – nécessaire au système cellulaire d'évacuation des déchets, et l'enzyme de synthèse du nicotinamide adénine dinucléotide (NAD), qui participe à de nombreuses réactions métaboliques. Quand un neurone est lésé chez une souris Wld^S, son axone dégénère bien plus lentement qu'il ne le ferait chez une souris saine.

Ce que raconte le cerveau

Comme il faudra encore attendre des années avant d'avoir des traitements efficaces de la sclérose latérale amyotrophique, on développe des appareils qui décodent des signaux du cerveau des patients paraplégés. Cela permet aux patients de communiquer, d'exécuter les fonctions de base d'un ordinateur et, dans certains cas, de commander des prothèses. Certaines de ces interfaces cerveau-machine fonctionnent grâce à des électrodes implantées chirurgicalement dans le cerveau et qui détectent l'activité d'un petit groupe de neurones dans le cortex moteur, le centre cérébral qui contrôle notamment les muscles.

Les interfaces cerveau-machine non invasives détectent l'activité électrique de millions de neurones *via* des électrodes posées sur le cuir chevelu du patient. Les neurobiologistes Jonathan Wolpaw, Theresa Vaughan, Eric Sellers et leurs collègues, du Département de la santé de l'État de New York, à Albany, ont développé ce type d'interface : elle détecte un signal cérébral lorsque quelque chose attire l'attention du patient. Dans ce système, les 72 lettres, chiffres, signes de ponctuation et de contrôle du clavier, disposés sur une grille de 17 colonnes par 17 lignes, s'allument selon une séquence rapide à l'écran de l'ordinateur pendant que le patient guette le symbole ou la fonction qu'il souhaite. Chaque fois que le symbole ou la fonction qui intéresse le patient s'allume, son cerveau émet une onde caractéristique : l'ordinateur analyse le décours temporel et d'autres caractéristiques de cette onde pour déterminer ce que la personne essaie de sélectionner.

Aujourd'hui, cinq patients atteints de la maladie de Charcot ont utilisé ce système pour écrire et communiquer. Un scientifique s'en sert même pour diriger son laboratoire de recherche. Un autre patient l'utilise pour transmettre des informations simples, mais essentielles au quotidien : « J'ai mal » ou « J'ai soif. »



3. In vitro, un motoneurone sain (en haut à gauche) possède des prolongements distincts, tandis qu'un motoneurone endommagé (en bas à gauche) a une forme tordue et des prolongements lésés. De plus, l'ADN dans le noyau du neurone malade forme des paquets (le cartouche). La plupart des malades meurent en quelques années à cause de la dégénérescence de ces neurones, mais le physicien Stephen Hawking (à droite) vit avec la maladie depuis plus de 40 ans.

En étudiant des neurones de souris en culture, l'équipe de J. Milbrandt a montré que la mutation des souris Wld^S augmente la quantité de NAD, ce qui protège les neurones. Par quel mécanisme ? Des concentrations élevées de NAD stimulent, dans la cellule, une voie biochimique qui serait impliquée dans la longévité accrue des ascarides (des parasites) et des drosophiles (des mouches). Qui plus est, d'autres petites molécules, tel le resvératrol que l'on trouve dans la peau des raisins noirs, stimulent cette voie de la longévité et protégeraient aussi les neurones lésés. Ainsi, plusieurs laboratoires pharmaceutiques tentent de développer des médicaments susceptibles de combattre la sclérose latérale amyotrophique et d'autres troubles neurodégénératifs en stimulant la voie biochimique mise en jeu par le NAD.

L'équipe de Peter Carmeliet, de l'Université catholique de Louvain, en Belgique, a opté pour un autre axe de recherche... par hasard. Elle a créé une souris transgénique incapable de produire le facteur de croissance de l'endothélium vasculaire (VEGF pour *vascular endothelial growth factor*), une protéine impliquée dans le contrôle de la croissance et de la perméabilité des vaisseaux sanguins. Étonnamment, ces souris développent une maladie des motoneurones.

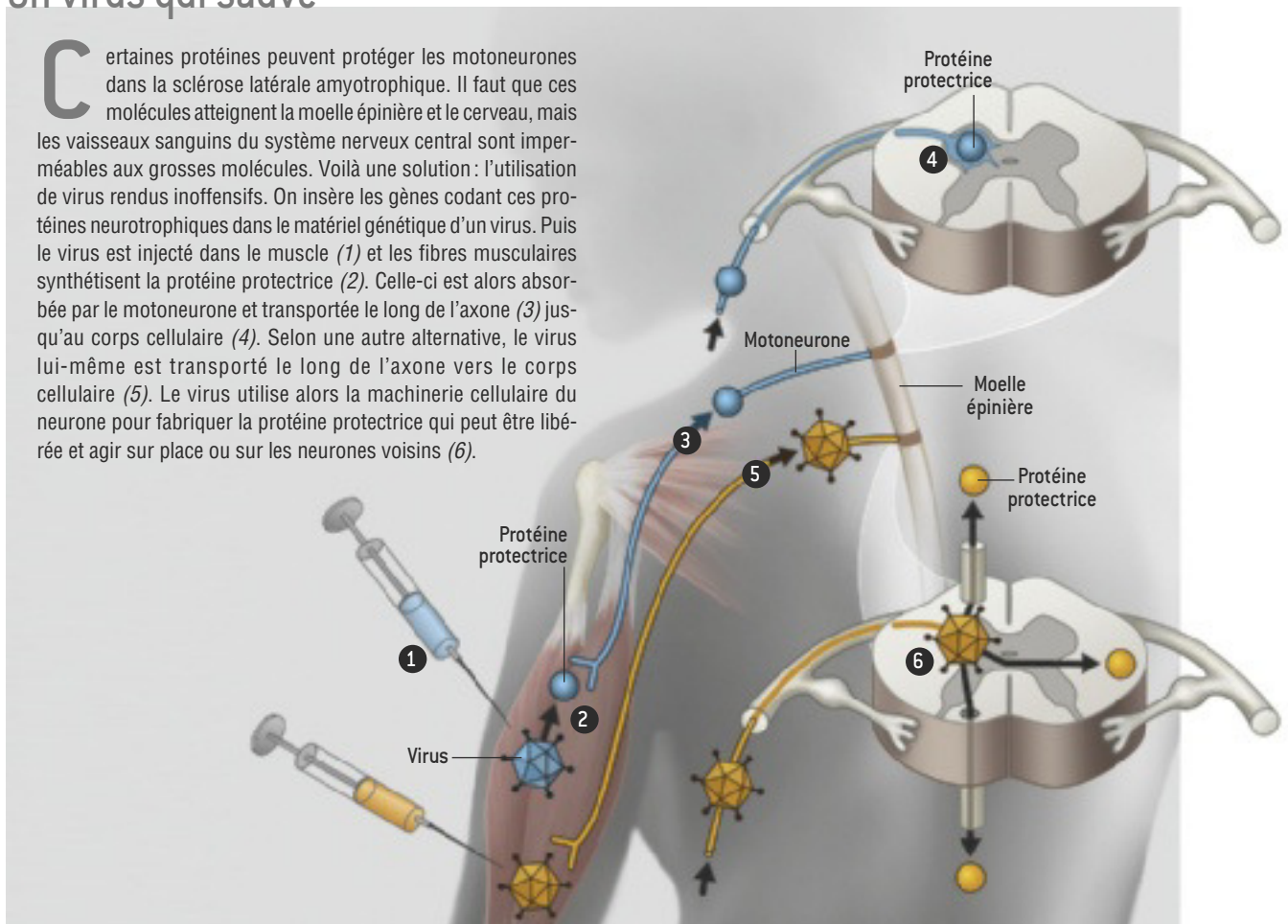
Les neurobiologistes ont alors tenté de ralentir la dégénérescence des neurones en leur apportant du VEGF. Mais administrer des protéines thérapeutiques dans le cerveau et la moelle épinière est compliqué parce que les parois des vaisseaux sanguins du système nerveux central sont imperméables aux grosses molécules. C'est alors que l'équipe de Mimoun Azzouz et Nicholas Mazarakis, de la Société britannique *Oxford Biomedica*, a développé une thérapie dite génique : elle a fabriqué un virus – inoffensif – qui produit le VEGF après avoir pénétré dans une cellule. Les chercheurs ont injecté le virus dans différents muscles des souris déficientes en VEGF ; les terminaisons des motoneurones, qui se connectent aux muscles, ont capturé le virus qui a ensuite été transporté aux corps cellulaires dans la moelle épinière (voir l'encadré page 59). Dans les motoneurones, le virus a produit du VEGF en quantité suffisante pour retarder le déclenchement de la sclérose latérale amyotrophique chez ces souris et en ralentir la progression.

Grâce à une autre technique d'administration de médicaments, P. Carmeliet et ses collègues ont utilisé une petite pompe pour transporter le VEGF dans le liquide cébrospinal – qui baigne le système nerveux central – de rats souffrant de sclérose latérale amyotrophique. La dégénérescence des neurones moteurs a encore été ralentie : l'administration du facteur de croissance protégerait le système nerveux de la maladie. Qui plus est, les patients atteints de sclérose latérale amyotrophique ont des concentrations sanguines en VEGF inférieures à la normale. Avec *NeuroNova*, un laboratoire pharmaceutique suédois, l'équipe essaie actuellement de développer des essais cliniques pour le traitement par le VEGF.

Une autre protéine qui protège les motoneurones est le facteur de croissance insulinoïque de type 1 (IGF-1, pour *insulinlike growth factor*). Fred Gage et ses collègues, de l'Institut *Salk* pour les études biologiques, à San Diego, ont construit un virus qui produit de l'IGF-1, puis ils l'ont injecté dans les muscles de souris dont le gène codant SOD1 est muté. Le virus a pénétré dans les terminaisons des motoneurones,

Un virus qui sauve

Certaines protéines peuvent protéger les motoneurones dans la sclérose latérale amyotrophique. Il faut que ces molécules atteignent la moelle épinière et le cerveau, mais les vaisseaux sanguins du système nerveux central sont imperméables aux grosses molécules. Voilà une solution : l'utilisation de virus rendus inoffensifs. On insère les gènes codant ces protéines neurotrophiques dans le matériel génétique d'un virus. Puis le virus est injecté dans le muscle (1) et les fibres musculaires synthétisent la protéine protectrice (2). Celle-ci est alors absorbée par le motoneurone et transportée le long de l'axone (3) jusqu'au corps cellulaire (4). Selon une autre alternative, le virus lui-même est transporté le long de l'axone vers le corps cellulaire (5). Le virus utilise alors la machinerie cellulaire du neurone pour fabriquer la protéine protectrice qui peut être libérée et agir sur place ou sur les neurones voisins (6).



puis a atteint les corps cellulaires, où il a produit l'IGF-1. Le traitement a permis une augmentation de 30 pour cent de l'espérance de vie des souris, probablement en protégeant les motoneurones et les cellules voisines. Les chercheurs préparent aujourd'hui des essais cliniques pour tester l'efficacité de l'IGF-1 chez des patients atteints de sclérose latérale amyotrophique.

De l'exercice pour protéger les motoneurones

Une des découvertes les plus surprenantes de ces dernières années est que l'exercice physique régulier stimulerait la croissance de nouveaux neurones, améliorerait l'apprentissage et augmenterait les concentrations de facteurs de croissance dans le système nerveux. En outre, dans des modèles animaux, on a montré que l'exercice physique protégerait les neurones après un traumatisme ou une maladie. Par exemple, F. Gage et ses collègues ont découvert qu'en faisant faire de l'exercice à leurs souris SOD1 dans des roues, leur durée de vie moyenne augmente de 120 à 150 jours. L'exercice stimulerait d'une façon ou d'une autre la production d'IGF-1 chez les souris, ce qui améliorerait leurs réactions motrices et protégerait même des neurones endommagés. Pour confirmer cette hypothèse,

les chercheurs ont combiné exercice physique et traitement à l'IGF-1 : les souris ont survécu encore plus longtemps (en moyenne 202 jours) que si elles ne subissaient que l'un ou l'autre des traitements.

Les thérapies avec les cellules souches – des cellules immatures qui peuvent se différencier en n'importe quelle cellule, notamment en neurones – permettraient aussi de traiter la sclérose latérale amyotrophique et d'autres maladies neurologiques. Au départ, on pensait seulement utiliser les cellules souches pour remplacer les neurones endommagés. Mais depuis quelques années, plusieurs études ont montré que les cellules souches greffées pourraient délivrer les facteurs de croissance nécessaires aux neurones lésés. Ainsi, il serait non seulement possible de protéger les neurones endommagés, mais aussi de stimuler la régénération dans la moelle épinière.

Dans des modèles murins de sclérose latérale amyotrophique, on a montré que les cellules souches greffées migrent dans les régions où les neurones meurent. Le tissu lésé émettrait des signaux moléculaires pour attirer les cellules souches. Et comme on sait que les neurones moteurs et les cellules gliales voisines sont impliqués dans la maladie, on pourrait greffer des cellules souches capables non seulement de devenir des motoneurones, mais aussi des cellules gliales.

Les techniques d'ARN interférence seraient aussi un outil efficace. Ce mécanisme met en jeu de courts brins

Comme on a découvert des mutations dans le gène *sod1* associées à des cas familiaux de sclérose latérale amyotrophique, on a développé des souris génétiquement modifiées pour exprimer ce gène muté. Ces souris développent les mêmes caractéristiques pathologiques que celles de la maladie humaine, et elles ont permis de mieux comprendre les rouages moléculaires qui engendrent la dégénérescence des motoneurones.

L'une des questions fondamentales était de comprendre pourquoi seuls les motoneurones sont touchés dans la sclérose latérale amyotrophique, alors que la protéine mutée est exprimée dans toutes les cellules. Nous avons alors supposé que seuls les motoneurones seraient capables d'enclencher un programme de mort spécifique en réaction aux dégâts engendrés par la protéine *SOD1* mutée. Restait à trouver le mécanisme impliqué.

Pour exécuter un programme de mort cellulaire, le système immunitaire dispose de plusieurs mécanismes. Le plus utilisé pour réguler son armada de lymphocytes tueurs fait appel à l'activation de récepteurs, dits récepteurs de mort, dont un représentant est le récepteur Fas. Ce dernier est présent à la surface de toutes les cellules. Quand il est activé par son ligand FasL, il commande aux cellules de déclencher leur propre mort. Or FasL est produit soit par les cellules cibles elles-mêmes, soit par des cellules environnantes.

Nous avons montré que les motoneurones expriment des récepteurs Fas à leur surface et que leur activation par FasL déclenche un programme de mort cellulaire qui implique une cascade d'événements intracellulaires spécifique des motoneurones.

Or les motoneurones de souris, isolés en culture et exprimant des formes mutées de *SOD1*, sont particulièrement sensibles à la mort provoquée par l'activation de Fas; la survie des motoneurones dépendrait donc beaucoup de leur environnement cellulaire (et de la production de FasL), un point très important dans le cas de la sclérose latérale amyotrophique.

Ces découvertes dans des modèles de souris sont très pertinentes, car d'autres équipes ont trouvé des signes d'activation de la voie de mort cellulaire mise en jeu par Fas chez les patients atteints de sclérose latérale amyotrophique. Bloquer l'activation des récepteurs Fas de mort cellulaire serait une option thérapeutique prometteuse: on utilisera pour cela la thérapie génique qui permet de délivrer efficacement aux motoneurones, *via* un transfert viral, des molécules interférant avec l'activation de la voie Fas.

On a récemment associé deux nouveaux gènes, *VAP-B* et *senataxin*, à des formes familiales de sclérose latérale amyotrophique et d'amyotrophie spinale infantile – une maladie neurologique héréditaire, dégénérative, qui touche les jeunes enfants, alors paralysés.

Le développement et l'étude de nouveaux modèles animaux et cellulaires exprimant les formes mutantes de *VAP-B* ou de *senataxin* permettront de déterminer les mécanismes communs aux différents modèles de sclérose latérale amyotrophique et donc d'identifier des candidats pertinents pour mettre au point de nouvelles thérapies.

Brigitte Pettmann et Cédric Raoul, Équipe INSERM-Avenir, Institut de biologie du développement de Marseille-Luminy

d'ARN qui se lient spécifiquement à certains ARN messagers – les molécules mobiles de l'information génétique qui permettent la synthèse des protéines; la liaison empêche les ARN messagers de synthétiser les protéines correspondantes. Ainsi, on infecte les cellules cibles avec des virus qui codent des brins d'ARN interférents, ces derniers empêchant alors la cellule de fabriquer des protéines toxiques. Une équipe dirigée par l'un d'entre nous, P. Aebischer, et M. Azzouz, a réduit au silence le gène muté codant *SOD1* en utilisant des ARN interférents: les souris transgéniques ont survécu plus longtemps.

De fait, d'autres scientifiques ont lancé un essai clinique pour utiliser les ARN interférents chez des patients atteints de la forme familiale de la sclérose latérale amyotrophique provoquée par une mutation du gène codant *SOD1*. Dans les premiers tests, une pompe délivrera les ARN interférents dans le liquide cébrospinal des patients. Ces molécules interagissent avec les ARN messagers avant que ces derniers n'assemblent les acides aminés pour fabriquer la protéine *SOD1* toxique dans les neurones et les cellules gliales. Si cette nouvelle technique donne de bons résultats, elle pourrait également être utilisée pour traiter d'autres maladies neurodégénératives dues à des gènes défectueux.

Mais toutes ces démarches thérapeutiques ont leurs limites. Pour que les médecins puissent administrer des facteurs de croissance ou des brins d'ARN interférents *via* des virus, ils doivent d'abord s'assurer que la procédure est sans danger. De plus, les cliniciens doivent évaluer combien de muscles il faudra traiter pour obtenir une amélioration fonctionnelle.

Dans l'idéal, il faudrait concevoir des combinaisons de plusieurs traitements – facteurs de croissance et ARN interférents par exemple – pour augmenter davantage la survie des patients. Une association à but non lucratif, nommée *Association SLA*, réalise déjà des essais cliniques pour une combinaison de médicaments: le célécoxib (un anti-inflammatoire qui ralentit la destruction des neurones par les cellules gliales hyperactives) et la créatine (un acide aminé qui favorise le fonctionnement des mitochondries). Ces deux molécules pouvant atteindre les motoneurones ont déjà prolongé la vie d'animaux souffrant de sclérose latérale amyotrophique.

Bien entendu, il serait préférable de prévenir la maladie de Charcot, plutôt que d'en traiter les symptômes. Cela ne sera envisageable que lorsqu'on aura identifié tous les facteurs (exercice physique, alimentation, gènes, etc.) impliqués dans la maladie...

Patrick AEBISCHER dirige l'École polytechnique fédérale de Lausanne, en Suisse, et **Ann KATO** est professeur de neurosciences à l'Université de Genève.

M. COLEMAN, *Axon degeneration mechanisms: commonality amid diversity*, in *Nature Reviews Neuroscience*, vol. 6, pp. 889-898, 2005.

G. RALPH et al., *Silencing mutant SOD1 using RNAi protects against neurodegeneration and extends survival in an ALS model*, in *Nature Medicine*, vol. 11, pp. 429-433, 2005.

L. BRUIJN et al., *Unraveling the mechanisms involved in motor neuron degeneration in ALS*, in *Annual Review of Neuroscience*, vol. 27, pp. 723-749, 2004.

DONNEZ DE L'ÉLAN À VOTRE ANNÉE !

Abonnez-vous à Pour la Science !

12 n^{os} – 54 €

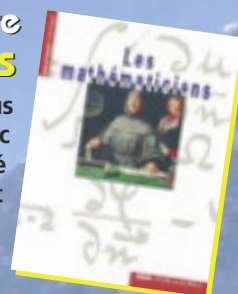


Photos non contractuelles

**POUR LA
SCIENCE**

En cadeau, le livre **Les Mathématiciens**

Cet ouvrage réunit les portraits des quinze plus grands mathématiciens. Leurs rapports avec leur époque, le regard qu'ils portent sur leur activité révèlent la genèse des idées nouvelles et mettent en lumière le caractère audacieux, dynamique et imaginaire de l'invention mathématique.



1996, 208 pages, Code Belin 001899, ISBN 2-9029-1899-2

BULLETIN D'ABONNEMENT

À retourner accompagné de votre règlement à :
Service Abonnements - Pour la Science
8 rue Férou - 75278 Paris Cedex 06

☐ **Oui, je m'abonne à Pour la Science seul : 1 an (12 n^{os}) et je reçois en cadeau le livre *Les Mathématiciens*.
54 € au lieu de 71,85 €* (prix de vente au numéro) Participation aux frais de port : 12 €/an pour l'Europe (hors France) et 25 €/an pour les autres pays.**

☐ J'accepte de recevoir par email des informations de Pour la Science.

☐ J'accepte de recevoir par email des informations des partenaires commerciaux de Pour la Science.

* prix en kiosque

Mes coordonnées :

Nom, prénom : _____

Adresse : _____

Ville : _____

Code postal : _____ Pays : _____

Téléphone* : _____

E-mail* : _____

Date de naissance* : _____

* facultatif

Je règle par :

☐ Chèque à l'ordre de Pour la Science

☐ Carte bancaire

Numéro de la carte

--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--	--

Date d'expiration

--	--	--	--

Signature obligatoire

En application de l'article 27 de la loi du 6 janvier 1978, les informations ci-contre sont indispensables au traitement de votre commande. Elles peuvent donner lieu à l'exercice du droit d'accès et de rectification auprès de Pour la Science. Par notre intermédiaire, vous pouvez être amené à recevoir des propositions d'autres organismes. En cas de refus de votre part, il vous suffit de nous prévenir par simple courrier.

PLS364

Liquides et glaces de spin

Dans certains réseaux cristallins, les énergies d'interaction entre aimants atomiques ou moléculaires ne peuvent être minimales toutes à la fois. Cette « frustration » conduit à des propriétés inhabituelles, que les physiciens cherchent à vérifier sur des matériaux réels. Avec quelques succès.

Rafik Ballou • Claudine Lacroix

Bien que l'homme ait une expérience des phénomènes magnétiques depuis plusieurs millénaires, la compréhension et l'utilisation du magnétisme sont longtemps restées rudimentaires. La richesse de ces phénomènes n'est apparue qu'aux XIX^e et XX^e siècles, avec la découverte, l'étude et l'exploitation de matériaux et de comportements magnétiques divers – des développements qui sont allés de pair avec les progrès réalisés dans la synthèse de matériaux nouveaux. Ce faisant, pour comprendre le magnétisme de la matière, les physiciens ont au cours des dernières décennies développé des concepts et des techniques qui s'appliquent à de nombreuses questions à caractère plus universel, telles l'étude des changements d'état de la matière ou l'étude d'un système de N corps en interaction.

L'exploration du magnétisme s'est ainsi révélée fructueuse à bien des égards. Elle n'est pas achevée, loin de là. En particulier, les physiciens ont récemment mis en évidence de nouveaux systèmes magnétiques, où intervient une « frustration » des interactions entre les éléments magnétiques (et microscopiques) du système. Nommés liquides de spin et glaces de spin (*voir la figure 1*) pour des raisons que nous expliquerons plus loin, ces systèmes présentent des propriétés originales et forcent les physiciens théoriciens à formuler de façon plus fine des notions fondamentales, tel le troisième principe de la thermodynamique. Les expérimentateurs, eux, cherchent à les mettre en évidence dans des matériaux réels. Et comme la notion de frustration apparaît dans d'autres situations, par exemple dans les empilements atomiques ou dans les cristaux liquides, l'intérêt dépasse le seul cadre du magnétisme.

Pour comprendre ce que sont les liquides et glaces de spin, commençons par rappeler le fondement microscopique des phénomènes magnétiques présentés par la matière. Cela nous permettra d'expliquer ce qu'est la frustration. Nous verrons alors comment, selon la géométrie du réseau formé par les éléments magnétiques du matériau, la frustration conduit à des propriétés originales, et nous donnerons quelques exemples de matériaux où ces propriétés interviendraient.

Les propriétés magnétiques de la matière résultent de celles de ses constituants, les atomes et les molécules, objets

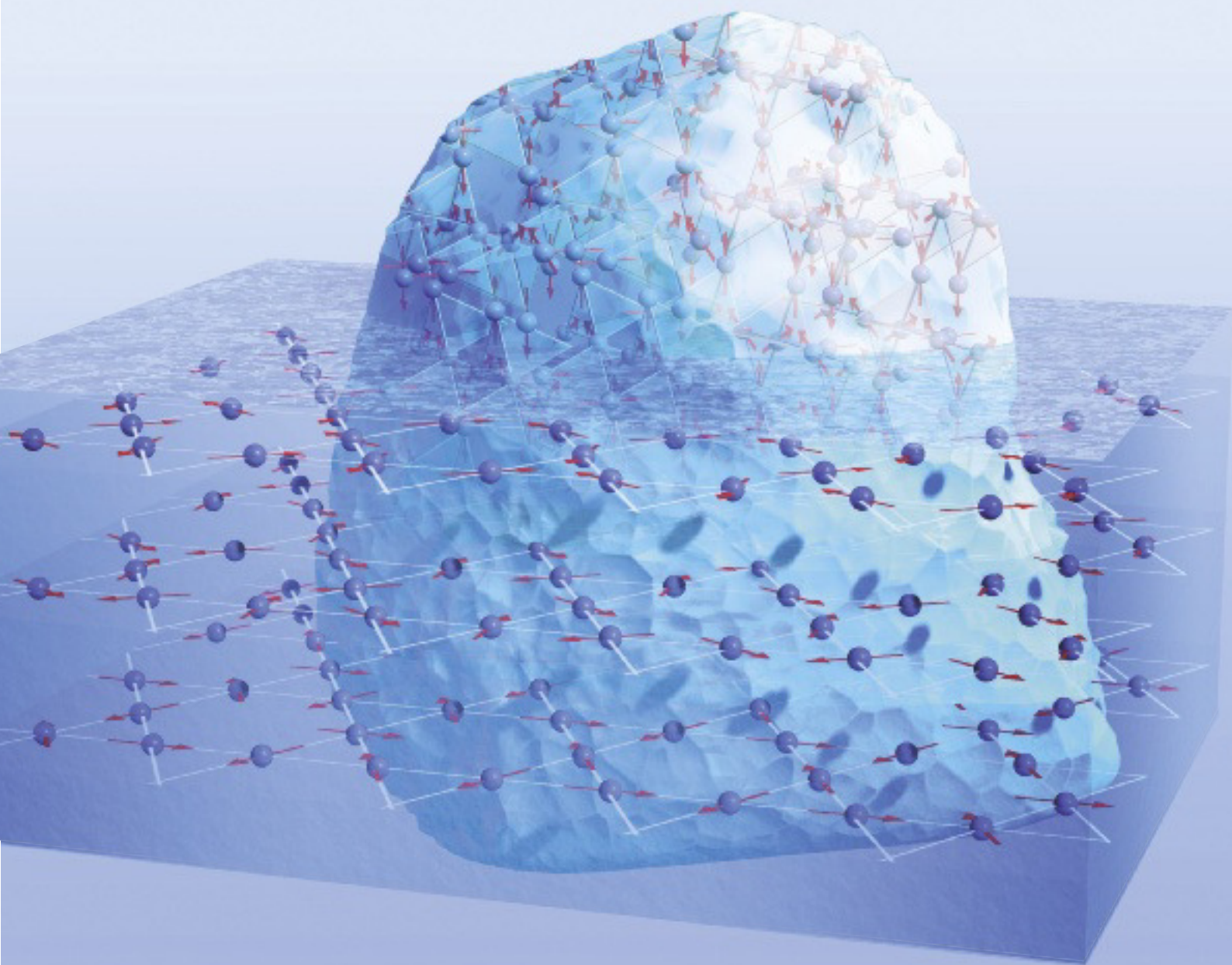
qui eux-mêmes sont composés de particules (électrons, protons, neutrons) dotées de caractéristiques magnétiques. Le rôle principal est dévolu aux électrons, ou plutôt aux cortèges qu'ils forment autour des noyaux atomiques : le magnétisme nucléaire, dû aux noyaux atomiques, est d'intensité beaucoup plus faible.

Chaque électron a un moment cinétique intrinsèque, le spin, grandeur quantique qui lui confère un moment magnétique. L'électron constitue ainsi une sorte d'aimant élémentaire. Dans un atome ou une molécule, les spins des électrons et les moments cinétiques associés à leur mouvement orbital se combinent selon des règles complexes pour donner un spin total. Si ce spin résultant est non nul, l'atome ou la molécule a un moment magnétique non nul. Par abus de langage, les spécialistes du magnétisme emploient souvent le terme de spin au lieu de moment magnétique, et c'est ce que nous ferons nous aussi.

Spins parallèles ou antiparallèles

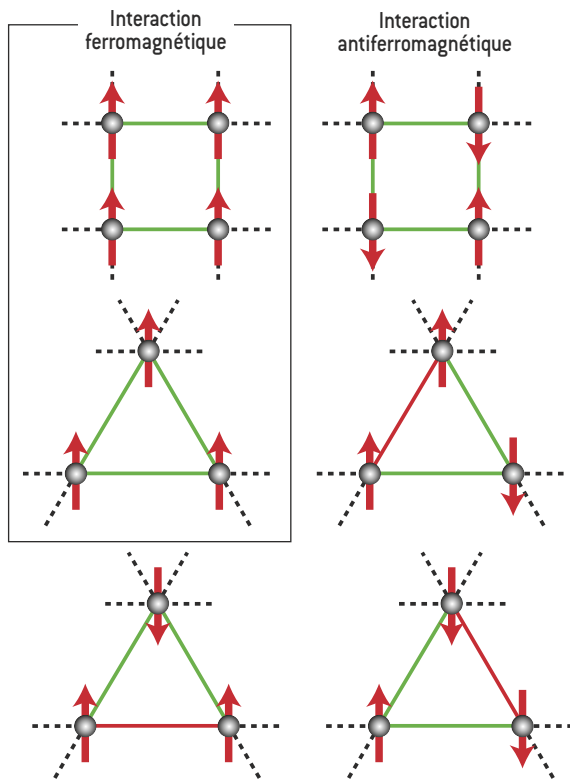
Les spins atomiques interagissent entre eux et avec un éventuel champ magnétique extérieur, et les propriétés magnétiques d'un matériau en découlent. Ainsi, dans une substance dite ferromagnétique, l'interaction des spins est telle que l'énergie d'une paire de spins voisins est minimale si les deux spins pointent dans une même direction (spins dits parallèles). Tout système physique ayant tendance à minimiser son énergie, il s'ensuit qu'à température suffisamment basse (au-dessous de 770 °C dans le cas du fer), une majorité de spins pointent dans la même direction : le matériau présente alors une aimantation spontanée macroscopique (*voir la figure 2*).

Le cas dit antiferromagnétique est plus subtil. Sa possibilité a été découverte théoriquement en 1936 par le physicien français Louis Néel (prix Nobel en 1970) et son existence a été confirmée en 1949 sur l'oxyde de manganèse (MnO) par les Américains Clifford Shull et Stuart Smart. Une interaction antiferromagnétique signifie que l'énergie d'interaction d'une paire de spins voisins est minimale lorsque les deux spins sont d'orientations opposées (spins dits antiparallèles). À suffisamment basse température (au-dessous de -151 °C pour l'oxyde de manganèse) et si la géométrie

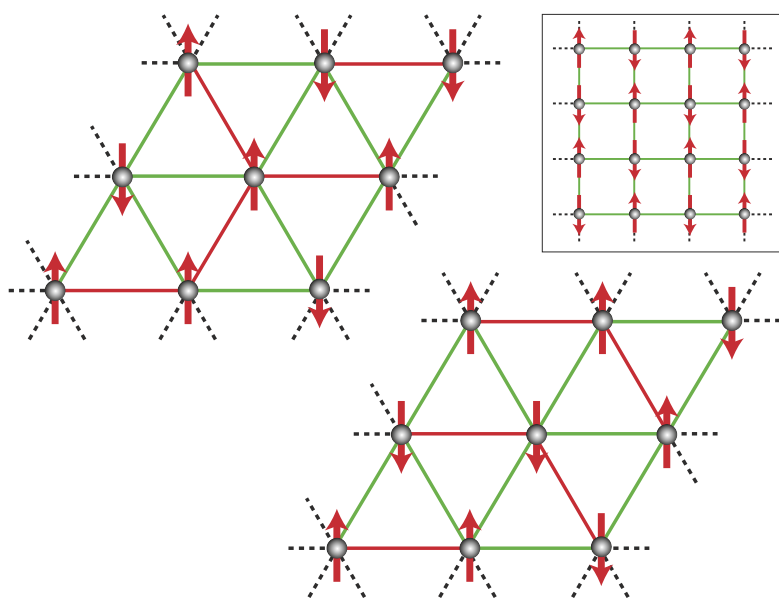


1. Les liquides et glaces de spin désignent des phases magnétiques pouvant s'établir dans des réseaux cristallins d'atomes ou molécules dotés d'un moment magnétique, ou spin, et disposés de telle façon que les spins ne peuvent s'orienter de manière à minimiser simultanément l'énergie d'interaction de toutes les paires de spins voisins. Cette frustration des interactions donne dans certains cas des « liquides de spin », où une infinité de configurations différentes correspondent

à la même énergie totale minimale [cas du réseau plan dit kagomé, à interactions antiferromagnétiques, illustré ici avec trois configurations planes de même énergie minimale]. La frustration peut aussi produire une « glace de spin », où les orientations des spins dessinent une structure équivalente à celle d'une des formes de la glace d'eau [cas du réseau dit pyrochlore, formé de tétraèdres, à interactions ferromagnétiques].



2. Pour une interaction ferromagnétique, deux spins voisins tendent à s'orienter dans le même sens [situation d'énergie minimale]. Pour une interaction antiferromagnétique, au contraire, l'énergie est minimale quand les deux spins sont antiparallèles. Pour quatre spins aux sommets d'un carré, les interactions peuvent être toutes satisfaites (*traits verts*), dans le cas ferromagnétique comme dans le cas antiferromagnétique. En revanche, pour trois spins aux sommets d'un triangle équilatéral, les interactions antiferromagnétiques ne peuvent être toutes satisfaites : au moins l'une d'elles (*traits rouges*) est frustrée.



3. Dans un réseau carré où les spins voisins ont des interactions antiferromagnétiques, il n'y a pas frustration : toutes les interactions sont satisfaites par une configuration où les orientations des spins sont alternées (*schéma encadré*). En revanche, dans un réseau triangulaire, il y a frustration puisque les interactions ne peuvent être toutes satisfaites simultanément. Ainsi, dans l'état d'énergie minimale, chaque triangle comporte une liaison frustrée (*traits rouges*). Un grand nombre de configurations différentes de spins correspondent à la même énergie minimale ; ici, deux configurations d'énergie minimale ont été représentées.

du réseau cristallin s'y prête, on a alors une configuration ordonnée où les spins orientés dans une direction alternent avec autant de spins pointant dans la direction opposée. À cette configuration correspond une aimantation nulle (voir l'exemple du réseau carré sur la figure 3).

Ainsi, dans l'exemple du réseau carré, comme dans le cas du réseau cristallin de l'oxyde de manganèse, toutes les interactions antiferromagnétiques sont « satisfaites », au sens où toutes les paires de spins voisins sont dans leur état antiparallèle, d'énergie minimale (on suppose ici qu'un spin n'interagit notablement qu'avec ses plus proches voisins). Il existe cependant des situations où toutes les interactions binaires ne peuvent être minimisées simultanément. On parle alors de frustration.

Une frustration peut être induite par un simple effet de géométrie. Pour l'illustrer, considérons le « modèle d'Ising », imaginé par le physicien allemand Wilhelm Lenz en 1920 et analysé par son étudiant Ernst Ising en 1925. Ce modèle simple, mais fécond et célèbre, suppose que les spins sont situés aux nœuds d'un réseau périodique dans le plan ou dans l'espace, et que chaque spin ne peut prendre que deux valeurs opposées, qu'on représente dans les dessins par des flèches vers le haut ou vers le bas. Le modèle d'Ising correspond aux substances dont les moments magnétiques sont contraints, par la structure anisotrope du matériau, à s'aligner selon un axe fixe. C'est le cas de nombreux composés de terres rares (famille de 16 éléments métalliques qui inclut les lanthanides et l'yttrium).

Un grand nombre d'états pour le même minimum d'énergie

Supposons que les interactions de spins voisins soient antiferromagnétiques et que le réseau soit triangulaire, c'est-à-dire formé de triangles équilatéraux se touchant par leurs côtés (voir les figures 2 et 3). Ce système présente une frustration géométrique : dans chaque triangle de trois spins proches voisins, il y a au moins une paire de spins frustrée (spins parallèles au lieu d'être antiparallèles). En revanche, si le réseau est carré ou cubique, il n'y a plus frustration : toutes les interactions des spins proches voisins peuvent être simultanément satisfaites, et l'énergie totale atteint alors un minimum profond. Dans cette situation d'énergie minimale, une fois choisie l'orientation d'un spin, celle de tous les autres est fixée de façon univoque. On dit dans ce cas qu'il existe un « ordre magnétique ». Remarquons qu'ici, lorsque l'énergie totale est minimale, on a deux configurations possibles de spins, l'une étant obtenue de l'autre par le renversement de tous les spins.

À l'inverse, en situation de frustration, il existe une multiplicité de configurations possibles pour l'énergie totale la plus basse. Ainsi, dans le cas de trois spins placés aux sommets d'un triangle, on a six configurations possibles de même énergie minimale. Il est ici possible d'inverser un seul spin sans changer l'énergie du trio, et il n'y a donc pas d'ordre magnétique. Quand plusieurs configurations ou états correspondent à une même énergie, les physiciens parlent de dégénérescence. Dans le cas de notre trio de spins, on a ainsi une dégénérescence non triviale de l'état d'éner-

gie minimale du système (la dégénérescence triviale correspond au renversement de tous les spins).

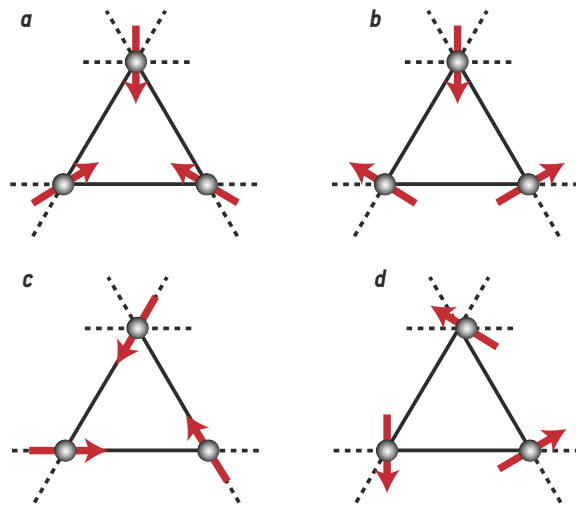
Pourquoi insister sur la notion de dégénérescence ? Parce qu'elle est liée à l'entropie, l'une des grandeurs fondamentales de la thermodynamique. À énergie donnée, l'entropie est une grandeur proportionnelle au logarithme du nombre d'états accessibles au système pour cette énergie. Autrement dit, plus le nombre d'états *a priori* possibles pour le système est grand, plus l'entropie est élevée. Or les propriétés thermodynamiques du système, telles que la chaleur spécifique (la quantité de chaleur qu'il faut fournir pour élever de un degré la température, par unité de quantité de matière), dépendent de l'entropie. Ainsi, une dégénérescence importante de l'état d'énergie minimale d'un système se répercutera sur son comportement thermodynamique à très basse température, le zéro absolu ($-273,15^\circ\text{C}$) étant la température où tout système se met dans son état d'énergie minimale, nommé état fondamental.

Les concepts de la thermodynamique n'ayant de sens que pour des systèmes à très grand nombre de constituants, l'important n'est pas la dégénérescence pour un trio de spins, mais la dégénérescence pour un grand réseau, comportant une multitude de spins. Or en présence de frustration, la dégénérescence augmente parfois exponentiellement avec le nombre de spins, ce qui correspond à une entropie non nulle au zéro absolu. Une telle situation indique que l'énoncé habituel du troisième principe de la thermodynamique, selon lequel l'entropie de tout système est nulle au zéro absolu, doit être nuancé.

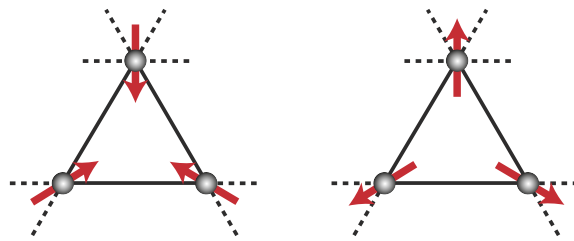
Cette conséquence de la frustration avait été perçue dès 1950 par le physicien suisse Grégory Wannier et par le physicien néerlandais Raymond Houtappel, qui ont étudié le modèle d'Ising dans le cas d'un réseau triangulaire plan. On ne connaissait à l'époque qu'un seul cas de système ayant une entropie non nulle à température absolue nulle. Étudié par Linus Pauling en 1935, il s'agit des configurations de position des atomes d'hydrogène dans une des variétés cristallines de la glace d'eau, système où une frustration géométrique est également présente. Nous en verrons un analogue magnétique plus loin, la glace de spin.

Considérons à nouveau un triangle formé de trois spins identiques, assimilés cette fois à des vecteurs $\vec{S}_1$, $\vec{S}_2$ et $\vec{S}_3$ pouvant *a priori* pointer dans n'importe quelle direction (et non à des entités ne pouvant prendre que deux valeurs opposées le long d'un axe). On parle dans ce cas de « spins de Heisenberg », par opposition à des « spins d'Ising ». On montre que pour un triangle, l'énergie totale (à une constante additive près) due à l'interaction des spins est proportionnelle à Σ^2 , où Σ est le module de la somme vectorielle $\vec{S}_1 + \vec{S}_2 + \vec{S}_3$; le coefficient de proportionnalité est positif pour des interactions antiferromagnétiques et négatif pour des interactions ferromagnétiques.

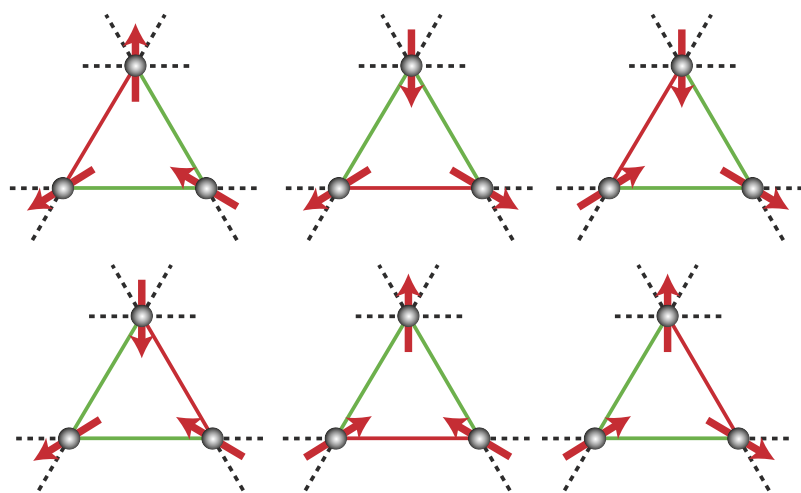
Ainsi, pour des interactions antiferromagnétiques, cette énergie atteint son minimum si Σ est minimal. Lorsque les spins sont libres de tourner dans toutes les directions, ce minimum d'énergie est atteint pour toutes les configurations où les spins font entre eux des angles de 120 degrés et où Σ est nul. Si les spins restent dans le plan du triangle, on trouve deux configurations de base correspondantes (voir la figure 4). Toutes les autres configurations de l'état fondamental



4. Si les spins sont libres de tourner, l'état d'énergie minimale d'un triangle à interactions antiferromagnétiques s'obtient quand la somme vectorielle des spins est nulle (angles de 120 degrés pour chaque paire de spins). On a deux configurations de base, qui sont planes (a et b). Toutes les autres s'en déduisent. Ainsi, les configurations planes c et d se déduisent de la configuration a par une rotation (d'angles -30° et $+120^\circ$ respectivement) appliquée à chaque spin. Des configurations non planes s'obtiennent en faisant tourner deux spins d'un même angle autour d'un axe parallèle au troisième spin.



5. Dans certains réseaux triangulaires, des anisotropies contraignent les spins à s'aligner selon les axes de symétrie. L'énergie minimale sur un triangle avec interactions antiferromagnétiques est obtenue avec deux configurations de spins, inverses l'une de l'autre.



6. Si les spins sont contraints de s'aligner selon les axes de symétrie du triangle et si les interactions sont ferromagnétiques, on a six configurations distinctes pour l'état d'énergie minimale. Dans chacune de ces configurations, on a deux liaisons plutôt satisfaites (angle aigu, traits verts) et une liaison plutôt frustrée (angle obtus, traits rouges).

s'obtiennent en faisant tourner deux des spins d'un même angle α autour de l'axe formé par le troisième spin. L'angle α étant un paramètre que l'on peut modifier continûment, l'état fondamental a une dégénérescence continue. Elle est de nature différente de la dégénérescence vue précédemment pour les spins d'Ising sur un triangle.

Si l'anisotropie du matériau impose que les spins soient dirigés selon les axes de symétrie du triangle, seules deux configurations minimisent l'énergie (voir la figure 5). En revanche, si les interactions sont ferromagnétiques, l'énergie est minimale quand Σ est maximal, et on a alors six configurations possibles pour l'état fondamental (voir la figure 6).

Que devient l'état fondamental pour tout un réseau de spins de Heisenberg constitué de triangles ? Examinons le cas antiferromagnétique sur un réseau triangulaire, où chaque spin est commun à six triangles (voir la figure 7), et sur un réseau kagomé (voir la figure 8), où chaque spin n'appartient qu'à deux triangles (kagomé est, au Japon, le nom d'un motif répandu de tressage de paniers). D'après ce que nous avons vu plus haut, si les spins sont libres de tourner, l'énergie est minimale lorsque la somme Σ des spins est nulle sur chaque triangle.

Prenons un triangle formé de trois spins orientés à 120 degrés les uns par rapport aux autres, configuration qui minimise l'énergie du triangle. Essayons de satisfaire les interactions de proche en proche, en commençant par les triangles adjacents. On constate que sur le réseau triangulaire, la configuration de tous les spins est complètement déterminée par la configuration du triangle initial. Le réseau triangulaire antiferromagnétique présente ainsi un ordre magnétique.

Il en va tout autrement avec le réseau kagomé antiferromagnétique. Dans un triangle adjacent, qui a un sommet commun avec le triangle initial, la configuration de celui-ci ne détermine que l'état d'un seul spin. Ce triangle adjacent possède une infinité continue, paramétrée par l'angle α défini précédemment, de configurations à 120 degrés avec un spin fixé et deux spins libres (voir la figure 8). En procédant de proche en proche, on voit ainsi qu'il existe une infinité de configurations possibles pour le réseau kagomé, qui toutes correspondent à la même énergie minimale.

On voit donc une différence essentielle entre les deux réseaux. Dans le réseau triangulaire, où chaque spin appartient à six triangles, la « connectivité » est grande et impose

un ordre magnétique. Dans le réseau kagomé, où chaque spin n'appartient qu'à deux triangles, la connectivité est faible et autorise un grand nombre (une infinité) d'états de même énergie ; le nombre de configurations d'énergie minimale croît de façon exponentielle avec le nombre de triangles, même si on se limite aux configurations planaires.

Les réseaux que nous avons examinés jusqu'ici sont bidimensionnels, mais on peut aussi concevoir des réseaux tridimensionnels à frustration géométrique.

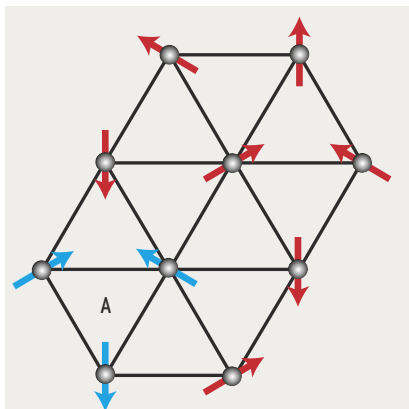
La cellule élémentaire à considérer est alors le tétraèdre, au lieu du triangle. Quatre spins situés aux sommets d'un tétraèdre forment quatre triangles et on montre alors que si les interactions sont antiferromagnétiques, l'énergie totale est minimisée lorsque la somme des quatre spins de chaque tétraèdre est nulle. En assemblant des tétraèdres par paires ayant un sommet commun, on obtient l'équivalent tridimensionnel du réseau kagomé, qu'on appelle réseau pyrochlore (voir les figures 1 et 10). Pour des spins sans anisotropie, on a comme précédemment une dégénérescence continue.

Quand la température diminue jusqu'au zéro absolu, un tel système dont l'état fondamental est fortement dégénéré ne se fige pas dans un état unique, correspondant à un arrangement bien défini (périodique par exemple) des moments magnétiques. Ainsi, les systèmes frustrés que sont les réseaux kagomé ou pyrochlore n'ont pas, à basse température, un ordre magnétique fixé.

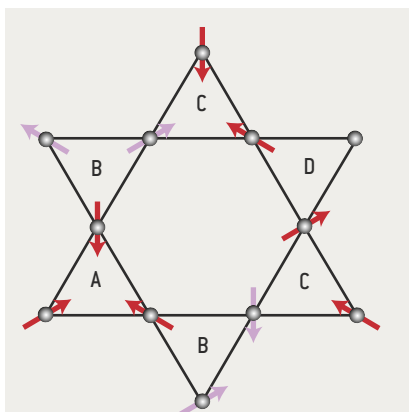
Souvent, l'ordre magnétique règne

Cependant, dans les matériaux réels, de nombreuses et diverses interactions plus faibles lèvent les dégénérescences et établissent un ordre magnétique. Il peut s'agir d'interactions entre spins non proches voisins, ou d'anisotropies cristallines qui favorisent certaines orientations des spins, ou encore d'interactions magnétoélastiques qui déforment le réseau cristallin (ce qui

modifie les distances entre spins, donc leurs interactions). C'est pourquoi il s'établit souvent un ordre magnétique complexe, mais à une température basse, caractéristique des énergies associées à ces interactions moins intenses et non aux interactions principales entre spins proches voisins. Parfois, cependant, les différentes interactions secondaires ne semblent pas toujours présentes ou suffisantes pour lever la dégénérescence. Mais on ne connaît encore que peu de matériaux réels dans lesquels les effets de frus-



7. Dans un réseau triangulaire à interactions antiferromagnétiques, si l'on se fixe une configuration d'énergie minimale pour un premier triangle [A], celles de tous les autres triangles sont fixées de proche en proche : les spins sont déterminés successivement en imposant que leur somme vectorielle soit nulle sur chaque triangle.



8. Dans un réseau kagomé à interactions antiferromagnétiques, la configuration choisie pour un premier triangle A laisse une grande liberté pour les autres. Même en contraignant les spins à rester dans le plan du réseau, on a deux choix possibles pour les triangles B, et deux choix à nouveau pour les triangles C [qui fixeront le triangle D].

tration géométrique dominant le comportement jusqu'aux plus basses températures. De tels matériaux sont activement recherchés.

En 1997 Mark Harris, Steven Bramwell et leurs collègues (une équipe de physiciens anglais, danois et allemands), ont découvert un nouvel état magnétique dans des systèmes pyrochlores, plus précisément dans le composé $\text{Ho}_2\text{Ti}_2\text{O}_7$: la « glace de spin ». Cet état, qui a aussi été observé par la suite sur le composé $\text{Dy}_2\text{Ti}_2\text{O}_7$, est induit par la frustration d'interactions ferromagnétiques et non antiferromagnétiques, ce qui est inhabituel. Dans ces composés, les atomes de dysprosium (Dy) ou de holmium (Ho) sont les seuls porteurs de moment magnétique, et ils forment un réseau pyrochlore de tétraèdres en contact par leurs sommets.

Une très forte anisotropie magnétocristalline gèle la direction de chaque spin selon l'axe de symétrie du tétraèdre en ce sommet. L'interaction étant ferromagnétique, les spins tendraient à s'orienter dans le même sens. Cependant, l'anisotropie rend cela impossible et, pour chaque tétraèdre, on a six états fondamentaux pour l'énergie la plus basse (voir la figure 9), où deux spins du tétraèdre pointent vers l'intérieur et les deux autres vers l'extérieur. Dans un tel état, quatre liaisons sont satisfaites (autrement dit, compte tenu de l'anisotropie, les deux spins de la liaison forment un angle aigu) et deux sont frustrées (angle obtus).

Quand les spins imitent la glace

Ces composés étaient les premiers systèmes magnétiques présentant une analogie forte avec un système frustré qui avait été identifié dans les années 1930, une forme cristalline de la glace d'eau notée I_h (voir la figure 10). Dans la glace I_h , les atomes d'hydrogène forment un réseau de tétraèdres similaire au réseau pyrochlore, les atomes d'oxygène étant situés à l'intérieur des tétraèdres. Cependant, l'atome d'oxygène n'est pas exactement au centre du tétraèdre : il est plus proche de deux des atomes d'hydrogène (et plus loin des deux autres) de façon à former une molécule H_2O . En dessinant sur chaque atome d'hydrogène une flèche dirigée vers l'atome d'oxygène le plus proche, l'état de chaque tétraèdre est tel qu'on a deux flèches vers l'intérieur et deux flèches vers l'extérieur, d'où une parfaite analogie avec les systèmes de spins ferromagnétiques des composés $\text{Dy}_2\text{Ti}_2\text{O}_7$ et $\text{Ho}_2\text{Ti}_2\text{O}_7$, ce qui explique aussi le choix du terme « glace de spin ».

Pour la glace d'eau I_h , Pauling avait calculé en 1935 la dégénérescence de l'état d'énergie minimale. Il s'ensuivait que l'entropie par site, au zéro absolu, était non nulle et égale à $(1/2) R \ln(3/2)$ (R étant la constante des gaz parfaits). Cette valeur est proche de celle mesurée sur les composés $\text{Dy}_2\text{Ti}_2\text{O}_7$ et $\text{Ho}_2\text{Ti}_2\text{O}_7$. Ce sont des cas uniques pour lesquels on a mesuré effectivement une entropie non nulle à très basse température. Cela semblait contredire le troisième principe de la thermodynamique, mais il n'en est rien : si l'on interprète correctement ce principe, celui-ci stipule seulement que l'entropie est une fonction croissante monotone de la température (ce qui est bien le cas ici).

Nous avons vu que dans les réseaux kagomé ou pyrochlore, les interactions antiferromagnétiques entre spins proches voisins ne peuvent à elles seules conduire à un ordre

magnétique stable. Néanmoins, elles imposent que la somme vectorielle des spins sur chaque cellule (triangle ou tétraèdre) soit nulle au zéro absolu. En conséquence, les spins voisins dans ces réseaux sont fortement corrélés et ne sont pas libres de fluctuer de façon indépendante à basse température. Par exemple, dans le réseau kagomé, chaque spin tend à tout instant à préserver un angle de 120 degrés avec ses quatre proches voisins.

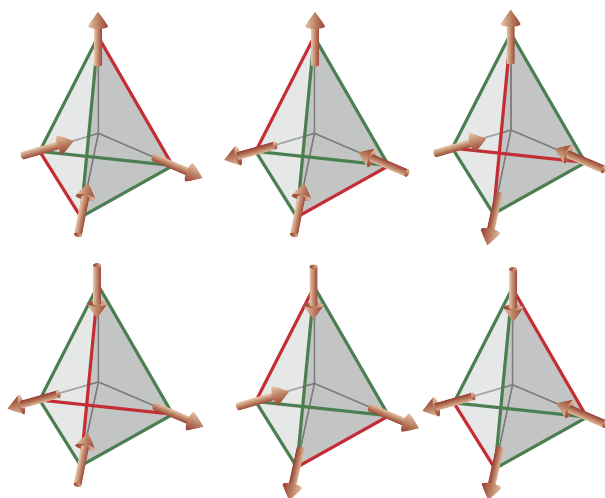
Des corrélations analogues à celles régnant dans un liquide

Ces corrélations entre spins sont analogues aux corrélations entre les positions des molécules dans un fluide, et c'est pourquoi on a nommé « liquide de spin » la phase magnétique caractérisant ces réseaux. Comme pour les corrélations de position dans les fluides, les corrélations de spin peuvent être mesurées par diffusion de neutrons sur le matériau. C'est ainsi qu'une phase de liquide de spin a été concrètement identifiée dans des composés tels que YMn_2 , à Grenoble en 1996 par Eddy Lelièvre-Berna, Bjorn Fak et l'un d'entre nous (Rafik Ballou), ou $\text{Tb}_2\text{Ti}_2\text{O}_7$, par l'équipe de John Greedan à l'Université McMaster, au Canada, en 1999.

Nous avons jusqu'ici décrit les spins comme des objets classiques que l'on peut assimiler à des vecteurs de longueur S , éventuellement contraints par une anisotropie (axiale ou planaire). Cette description s'applique à certains matériaux, mais pas à tous. En réalité, le spin est une grandeur quantique, obéissant à des lois particulières. Notamment, S ne peut prendre que des valeurs entières ou demi-entières (0, 1/2, 1, 3/2, 2, etc.) dans les unités appropriées.

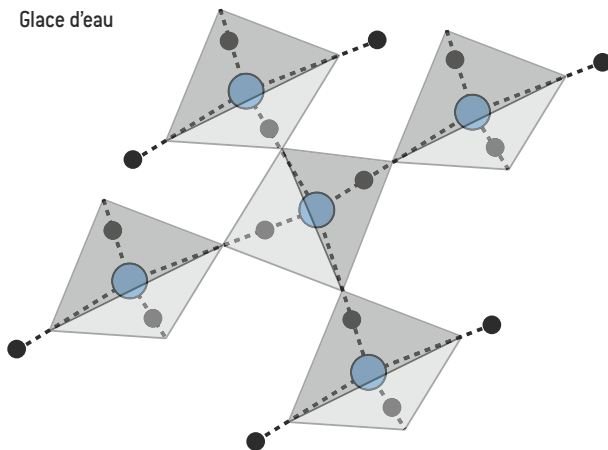
L'ordre par le désordre

Habituellement, un ordre magnétique dans un matériau est déstabilisé quand on introduit un désordre, par exemple en réalisant des substitutions chimiques qui diminuent la concentration de sites magnétiques. Certains systèmes à frustration géométrique présentent, contre toute attente, un effet inverse : alors que la frustration empêche le système de s'ordonner, un ordre est induit par augmentation du désordre. Ainsi, dans le minéral nommé jarosite de fer, un ordre magnétique bien défini apparaît quand on augmente la proportion d'ions non magnétiques dans le réseau kagomé formé par les ions de fer Fe^{3+} . Ici, le désordre chimique détruit localement la frustration géométrique : un triangle ne contenant que deux sites magnétiques n'est plus frustré. Notons qu'il existe également des processus de mise en ordre sous l'effet des fluctuations thermiques, mais ils sont plus subtils et donc moins aisés à mettre en évidence expérimentalement. Cela semble être le cas du composé pyrochlore $\text{Er}_2\text{Ti}_2\text{O}_7$, qui présente un ordre magnétique dans lequel tous les spins sont colinéaires à basse température. Ce phénomène de mise en ordre par un désordre est paradoxal, en ce sens que les fluctuations tendent généralement à détruire l'ordre (ce qu'elles font d'ailleurs si on élève trop fortement la température dans le cas du composé $\text{Er}_2\text{Ti}_2\text{O}_7$). Appelé ordre par le désordre, il fait partie de ces effets insolites propres à la frustration.

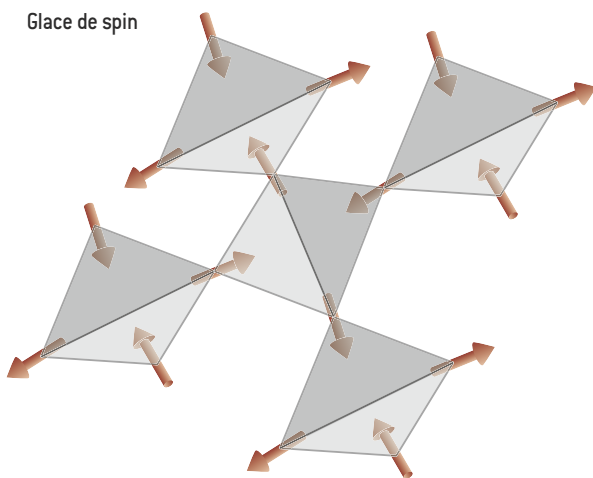


9. Pour un tétraèdre régulier de spins à interactions ferromagnétiques, et en supposant la direction des spins figée selon les axes de symétrie, on a six configurations de même énergie minimale. Il y a frustration : dans chacune de ces configurations, quatre liaisons sont satisfaites [angles aigus, traits verts] tandis que deux sont frustrées [angles obtus, traits rouges].

Glace d'eau



Glace de spin



10. La structure de la glace d'eau (H_2O) dans sa forme cristalline I_h , telle qu'elle est dessinée par les segments reliant les atomes d'oxygène (*en bleu*) aux atomes voisins d'hydrogène (*en noir*), est exactement de même type que celle dessinée par les spins d'un réseau pyrochlore à interactions ferromagnétiques dans l'état d'énergie minimale. C'est pourquoi ce système magnétique est nommé glace de spin.

La description classique n'est approximativement valable que lorsque S est assez grand. Plus S est petit, plus les aspects quantiques sont importants, et les analogies entre un liquide de spin et un liquide ordinaire ne fonctionnent alors plus.

En fait, pour des spins quantiques, les théoriciens ont établi que la frustration est à l'origine de propriétés inattendues et originales, qui ne sont pas encore toutes élucidées. Considérons tout d'abord deux spins $1/2$ en interaction antiferromagnétique. On montre que l'état fondamental de cette paire de spins est une certaine superposition quantique des deux configurations antiparallèles $(+1/2, -1/2)$ et $(-1/2, +1/2)$ possibles. Le spin total correspondant à cette superposition est nul ($S = 0$) : cet état, dit singulet, est non magnétique. Une certaine énergie le sépare de l'état excité immédiatement supérieur, de spin total $S = 1$.

On définit habituellement un liquide de spin quantique comme un système macroscopique, à grand nombre de constituants, dont l'état fondamental est de spin total nul ($S = 0$, où S est le spin de l'ensemble des constituants) et où le premier état excité magnétique, généralement un état $S = 1$, est d'énergie strictement supérieure à celle de l'état fondamental singulet. La différence des deux énergies, nommée gap de spin, est l'énergie minimale à fournir pour passer d'un état fondamental non magnétique (de spin nul) à un état magnétique (de spin total 1). L'existence d'un gap de spin non nul a des conséquences sur le comportement du système à très basse température. Par exemple, la susceptibilité magnétique (rapport entre l'aimantation du matériau et le champ magnétique appliqué) s'annule exponentiellement quand la température tend vers le zéro absolu, alors qu'en l'absence de gap de spin, elle varie comme une puissance de la température.

Un système de dimension deux ou trois peut-il être un liquide de spin quantique ? La question est longtemps restée sans réponse, car les calculs portant sur des spins quantiques sont généralement inextricables. En 1973, le physicien américain Philip Anderson (prix Nobel en 1977) avait suggéré que, grâce à la frustration, un liquide de spin quantique pouvait être réalisé par le réseau triangulaire de spins $1/2$.

La herbertsmithite, un liquide de spin quantique

Ph. Anderson avait proposé de décrire l'état fondamental de ce système bidimensionnel comme une superposition quantique d'états dont chacun correspond à un découpage du réseau en paires de spins voisins correspondant à des singulets de spin (*voir la figure 11*). Il a nommé RVB (pour *Resonating Valence Bond*, ou lien de valence résonnant) un tel état, qui est complexe puisqu'il existe un grand nombre de façons de décomposer le réseau triangulaire en paires de spins voisins.

En fait, la description RVB a été abandonnée pour le réseau triangulaire. Selon de nombreuses indications d'ordre théorique, ce réseau possède à température nulle un ordre antiferromagnétique à longue portée : ce n'est donc pas un liquide de spin. En revanche, d'après des études plus récentes, les états de type RVB conviennent bien à la description du

réseau kagomé de spins $1/2$: l'état fondamental se décrit bien comme une superposition quantique d'états formés de singulets à deux spins voisins, et c'est aussi le cas des états excités de basse énergie. L'état fondamental de ce réseau est bien un singulet (c'est-à-dire de spin $S = 0$), il existe un gap de spin séparant ce singulet du premier état de spin total $S = 1$, mais il existe tout un continuum d'états singulets, en fait des états RVB, dont les énergies sont intermédiaires entre l'énergie du fondamental et celle de l'état $S = 1$.

Existe-t-il un matériau correspondant au réseau kagomé de spins $1/2$ avec interactions antiferromagnétiques ? La réponse est oui : la herbertsmithite, un minéral rare de formule $\text{ZnCu}_3(\text{OH})_6\text{Cl}_2$. Les atomes de cuivre, de spin $1/2$, y forment des plans kagomé. En 2004, Daniel Nocera et ses collègues, à l'Institut de technologie du Massachusetts, ont réussi à synthétiser ce minéral, ce qui a ouvert la voie à l'étude d'échantillons suffisamment purs. Depuis 2006, plusieurs équipes ont ainsi publié des résultats montrant que la herbertsmithite ne présente pas d'ordre magnétique jusqu'à au moins 0,4 kelvin. Ce minéral est donc un candidat sérieux au titre de liquide de spin quantique ; cependant, les mesures des propriétés thermodynamiques à basse température, encore inachevées, devront le confirmer.

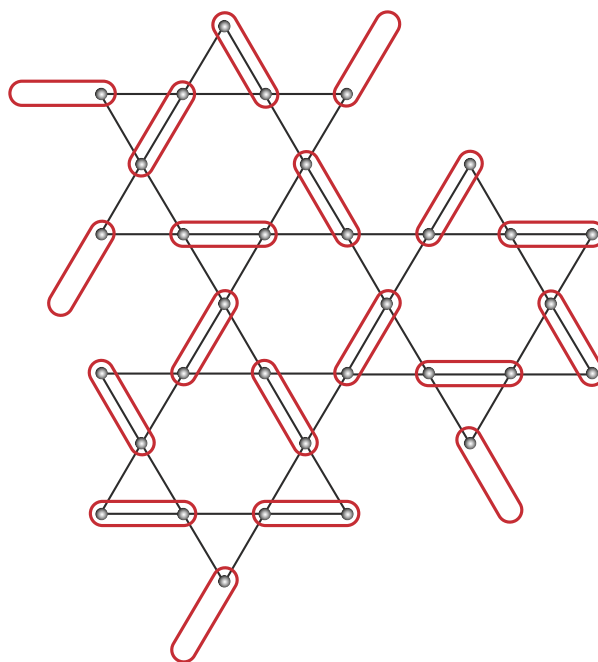
Les réseaux kagomé, donc la herbertsmithite, sont des systèmes bidimensionnels. Existe-t-il des exemples tridimensionnels de liquide de spin quantique ? En 1998, Benjamin Canals et l'un d'entre nous (Claudine Lacroix), à Grenoble, ont montré que le réseau pyrochlore de spins $1/2$ constitue un tel système. C'est le seul modèle théorique connu de liquide de spin quantique tridimensionnel. Ses propriétés sont sans doute semblables à celles du réseau kagomé. Cependant, aucun exemple expérimental de système pyrochlore quantique, restant désordonné jusqu'au zéro absolu, n'a été pour l'instant identifié, et les recherches se poursuivent.

De la frustration magnétique pour les réfrigérateurs ?

La frustration de spins pourrait intervenir dans d'autres contextes. Elle pourrait donner lieu à une phase de type verre de spin. Rappelons qu'un verre de spin est un système où les positions des spins, ou leurs interactions, sont désordonnées, ce qui conduit à des propriétés thermodynamiques particulières. Il semble que dans un système dépourvu de désordre, la frustration peut conduire à des propriétés analogues, par un mécanisme totalement différent. La phase correspondante, nommée verre de spin topologique et dont le composé nommé jarosite de fer est peut-être un exemple, est encore mal comprise.

Une autre question étudiée est celle de la supraconductivité dans les systèmes de spins frustrés. Ph. Anderson avait déjà proposé l'idée que l'état supraconducteur serait lié à un état de type RVB. Or plusieurs supraconducteurs à réseau pyrochlore de spins ont été découverts récemment, tels $\text{Cd}_2\text{Re}_2\text{O}_7$ ou KOs_2O_6 , et la compréhension de la supraconductivité dans ces systèmes fait l'objet de nombreux travaux.

Notons enfin que le phénomène de frustration pourrait conduire à des applications. En 2003, Mike Zhitomirsky, du CEA à Grenoble, a suggéré que les liquides de spin de



11. Le réseau kagomé à spins quantiques antiferromagnétiques constitue un liquide de spin quantique. Son état fondamental est décrit comme une superposition quantique d'états dont chacun correspond à une partition du réseau en paires de spins. Dans un tel découpage (dont un exemple est représenté ici), chaque paire de spins (entourée de rouge) est supposée être dans un état dit singulet, de spin total nul.

systèmes frustrés, tels que ceux que nous avons décrits, doivent présenter un effet magnétocalorique (changement de température causé par une variation de l'aimantation) géant.

Utilisé dans la réfrigération magnétique, l'effet magnétocalorique s'obtient lorsque, à une certaine température, on fait varier brutalement le champ magnétique appliqué au matériau sans échange de chaleur (désaimantation adiabatique). L'entropie doit alors rester constante, ce qui n'est possible que si la température varie (car l'entropie est une fonction du champ appliqué et de la température). Comme nous l'avons déjà mentionné, l'entropie des matériaux frustrés reste très grande à basse température, et l'on s'attend donc à un effet magnétocalorique supérieur de plusieurs ordres de grandeur à celui des matériaux utilisés habituellement pour la réfrigération magnétique. Or la recherche de matériaux utilisables pour la réfrigération magnétique connaît un grand essor, cette technique étant considérée comme plus écologique que les techniques classiques à base de gaz réfrigérants. On voit ici une utilisation potentielle inattendue des systèmes frustrés... auxquels seuls les théoriciens s'intéressaient initialement.

Rafik BALLOU et **Claudine LACROIX** sont physiciens à l'Institut Néel (CNRS et Université Joseph Fourier), à Grenoble.

B. GOSS LEVI, *New candidate emerges for a quantum spin liquid*, in *Physics Today*, vol. 60, n° 2, pp. 16-19, février 2007.

R. MOESSNER et A. P. RAMIREZ, *Geometrical frustration*, in *Physics Today*, vol. 59, n° 2, pp. 24-29, février 2006.

J. E. GREEDAN, *Frustrated rare earth magnetism : spin glasses, spin liquids and spin ices in pyrochlore oxides*, in *Journal of Alloys and Compounds*, vol. 408-412, pp. 444-455, 2006.

F. MILA, *Quantum spin liquids*, in *European Journal of Physics*, vol. 21, n° 6, pp. 499-510, 2000.

La redoutable efficacité des armes

Depuis le XIX^e siècle, archéologues et amateurs multiplient les expériences pour reproduire et tester les armements romains. Les résultats attestent d'un grand savoir-faire, produisant les armes défensives ou offensives qui ont fait la puissance de Rome.

Yann Le Bohec



1. L'artillerie jouait un grand rôle tant dans les batailles rangées que dans les sièges que livrait la légion romaine. Dans cette représentation d'artiste fondée sur nos connaissances actuelles, deux servants utilisent une chirobaliste, c'est-à-dire une grande arbalète sur pied, qui s'armait à l'aide d'un treuil.

romaines

Tout le monde sait en quoi consiste le métier de soldat : tuer. Un combattant doit savoir tuer l'ennemi sans se faire tuer lui-même ! Dans le cas de l'armée romaine, dont l'efficacité n'est plus à démontrer, cet art a atteint des sommets. D'où venait cette efficacité ? De plusieurs facteurs, comme la qualité de l'encadrement, la diversité des types d'unités, la pratique de l'exercice, l'art du siège et de la bataille.

Un autre facteur, essentiel, est d'autant plus intéressant à étudier qu'il a été négligé par le passé : l'armement. Certes, quelques érudits ont essayé de décrire les panoplies que montrent par exemple les bas-reliefs. Mais il est possible d'aller plus loin en recourant à l'archéologie expérimentale. Des particuliers ont ainsi pris l'habitude de s'habiller

en soldats romains, en recherchant la précision dans le détail. Souvent Anglo-Saxons et fils spirituels de l'*Ermine Street Guard*, une association britannique pionnière consacrée à cette activité, ils ont donné naissance à une discipline qu'ils appellent *re-enactment*, soit littéralement « représentation ». À leurs efforts s'ajoutent quelques travaux récents, où des chercheurs ou des amateurs essaient de reconstituer les armes romaines. Ce type d'études expérimentales apporte des indices restituant de façon réaliste le fonctionnement des armes, en l'occurrence des arcs et des flèches. Leur principale limite tient au fait que les Anciens procédaient de manière empirique. Les archéologues ayant rarement les moyens de recourir aux services de physiciens ou de chimistes, ils se bornent à reproduire des objets découverts en fouilles ou à les reconstituer le mieux possible.

350 000 hommes en armes, 150 000 légionnaires

Avant de nous pencher sur l'armement romain, rappelons l'ampleur de l'épopée militaire romaine. Au milieu du VIII^e siècle avant notre ère, la ville de Rome n'était qu'un petit village perdu au cœur du Latium. Ce n'est qu'au bout de plusieurs siècles de luttes que les légions se sont emparées de toute la région (en 336 avant notre ère). Puis l'impérialisme de la cité s'est développé : conquête de l'Italie, puis de tous les territoires bordant la Méditerranée. Au II^e siècle, le monde romain s'étendait de l'Écosse au Sahara, de l'Atlantique à la Mésopotamie ! Et le plus extraordinaire, c'est que 350 000 hommes seulement ont suffi pour maintenir cet immense empire intact pendant plus de quatre siècles, c'est-à-dire jusqu'à la prise de Rome par les Goths d'Alaric, en 410.

Évidemment, les soldats n'auraient pu, à eux seuls, assurer la cohésion d'un aussi vaste domaine. Il fallait l'adhésion des civils. Mais sans les militaires, rien n'aurait été possible. Qui étaient ces soldats ? L'armée romaine ayant beaucoup varié au cours des siècles, nous limiterons notre propos à la « belle époque » des I^{er} et II^e siècles de notre ère, époque d'extension maximale de l'Empire romain, pendant laquelle la formidable machine de guerre romaine comprenait plusieurs sortes d'unités.

L'élite de l'armée romaine était ce que l'on appelle « la garnison de Rome ». Les neuf cohortes de prétoriens, 500 fantassins chacune, formaient la garde impériale ; dans cette mission, elle était aidée par les 500 cavaliers gardes du corps. Les trois cohortes urbaines, aussi de 500 fantassins chacune, assuraient la police de la ville de Rome. Enfin, à leurs côtés, les pompiers de Rome, qui avaient été militarisés, étaient répartis en sept cohortes de 1 000 hommes.





Helmut Dinger, Heilmatherein, Verona



Romans in Britain

2. L'artillerie romaine avait de terribles effets, comme en attestent les restes de deux guerriers celtes tués lors d'un assaut de la légion en Grande-Bretagne. Tous les deux ont reçu une flèche de chirobaliste (une arbalète sur pied). L'un l'a reçue dans la colonne vertébrale. L'autre a été touché à la tête. On remarque la section carrée nette du trou dans le crâne, qui atteste de la vitesse d'impact, car aucune fissure n'est visible autour. *Ci-contre*, une autre arme d'artillerie romaine : le scorpion, une sorte de petite catapulte.



Lavard

3. Le savoir-faire guerrier des Romains n'était pas que dans les armes. Chaque soir, après la journée de marche, les légionnaires établissaient un camp quadrangulaire suivant un rituel immuable. Ils creusaient notamment quelque 2 500 mètres d'un fossé triangulaire profond de 2,25 mètres et large de 4,5 mètres. La terre était rejetée vers l'intérieur de façon à constituer un talus – un glacis – planté de pieux qui était en même temps un chemin de ronde. Le glacis qui entourait le camp ainsi fortifié était dans le même temps débroussaillé afin d'exclure toute attaque surprise. Pour effectuer de tels travaux d'Hercule, les légionnaires avaient besoin de bons outils.

Mais le principal instrument utilisé par les Romains pour assurer la défense de l'empire, c'étaient les légions et leurs auxiliaires. Leur armée comptait en effet de 25 à 35 unités de ce type (le nombre total a varié en fonction des créations, des dissolutions et des pertes au combat), chacune comptant environ 5 000 hommes. Une légion était divisée en 10 cohortes ou 30 manipules ou 59 centuries. Comme les prétoriens et les *urbanici* (des cohortes urbaines), les légionnaires étaient des soldats d'infanterie lourde, c'est-à-dire qu'ils possédaient une panoplie complète – casque, cuirasse et bouclier pour se protéger, glaive et javelot pour tuer. Avant la confrontation, des fantassins légers, mais aussi, comme l'a montré Michael Speidel, professeur d'histoire à l'Université d'Hawaï, les légionnaires, arrosaient l'ennemi de flèches ; des balles de plomb, lancées par des frondes, visaient à le harceler. Enfin, avant et pendant l'engagement, l'artillerie infligeait des pertes sévères.

L'affaire en vient aux triaires...

Au combat, les cohortes avançaient, séparées les unes des autres, sur trois lignes ; chacune comprenait un manipule d'« hastats » en première ligne et un manipule de « princes » en deuxième ligne. Ces soldats, formés pour l'assaut, étaient soutenus par un manipule de triaires, chargés d'arrêter la progression de l'ennemi en cas de grande difficulté (« L'affaire en vient aux triaires », un proverbe latin, voulait dire que la situation devenait très grave). Un solide encadrement, assuré par 59 centurions, cinq tribuns équestres, un tribun sénatorial et un légat, également sénatorial, permettait une bonne décentralisation du commandement pendant la bataille.

En tout temps, c'est-à-dire également en temps de paix, la légion était aidée dans sa tâche par des soldats dont le nom indique bien la subordination, puisqu'on les appelait les « auxiliaires » (de la légion). Ces derniers étaient à peu près aussi nombreux que les légionnaires. Ils fournissaient une cavalerie indispensable pour effectuer les liaisons et pour garantir la protection des flancs de l'infanterie lourde au combat. Ils donnaient également des fantassins légers, bien utiles contre des ennemis jugés faibles ou avant l'engagement contre des adversaires plus dangereux. On distinguait donc des

ails, qui regroupaient 500 ou 1 000 cavaliers, encadrées par un préfet ou un tribun et des décurions, et des cohortes, de 500 ou 1 000 fantassins également, encadrées elles aussi par un préfet ou un tribun et par six centurions. Des *numeri* sont attestés à partir du début du II^e siècle ; ces unités étaient recrutées en pays barbare et elles conservaient leurs caractères ethniques : les Syriens et les Thraces étaient réputés pour leur habileté à l'arc, les Gaulois et les Espagnols s'étaient rendus maîtres dans l'art de l'équitation, les Bretons donnaient de bons éclaireurs et les Germains des troupes de choc destinées à effrayer par leur cruauté.

La guerre, du 21 mars au 23 septembre...

Voilà pour ce que les historiens appellent « l'armée des frontières ». Reste la marine. Longtemps, on a jugé qu'elle était devenue inutile à partir du moment où Rome avait pris le contrôle de toutes les côtes et que, de toute façon, elle brillait par sa médiocrité. Dans un livre récent, Michel Reddé, un spécialiste d'archéologie militaire de l'École pratique des hautes études à Paris, a montré que les Romains avaient appris à construire les plus gros navires que l'Antiquité a connus, les plus solides et les mieux armés. De plus, la marine romaine assurait la logistique des opérations terrestres, surveillait les cours d'eau et les mers, et elle se préparait à repousser un ennemi encore inconnu : le bon stratège est celui qui prévoit tout, surtout l'imprévisible.

Pour des raisons religieuses et pratiques, la guerre durait du 21 mars au 23 septembre. Les dieux en avaient décidé ainsi, et il ne fallait s'aventurer chez l'ennemi qu'à partir du moment où le blé avait mûri, car une armée en marche se nourrissait en pillant les récoltes et les provisions de l'ennemi. Quand il ne se battait pas, le soldat s'entraînait. En ce temps-là, l'exercice, permanent et obligatoire, n'était pratiqué que par deux armées au monde : l'armée chinoise et l'armée romaine !

4. La panoplie du légionnaire romain était celle d'un soldat d'infanterie lourde, solidement appuyé au sol et entraîné à mener en groupe des assauts lents, mais puissants. Le légionnaire portait sur lui environ 25 kilogrammes, dont un casque, une cuirasse faite de bandes métalliques articulées (ou une cotte de mailles), un glaive (*ci-dessous*), un poignard, un bouclier renforcé de métal sur sa tranche supérieure (pour frapper au menton), etc. Grâce à cet équipement, une cohorte de légionnaires était une machine à tuer, lente mais efficace.



Au combat comme à l'entraînement, les soldats romains utilisaient un armement complexe, dont les reconstitutions expérimentales permettent de préciser le fonctionnement et l'efficacité. Après Napoléon III, l'empereur d'Allemagne Guillaume II fut sans doute l'un des tout premiers archéologues expérimentaux. Il avait lu *La chirobaliste*, un traité de Héron d'Alexandrie (I^{er} siècle). Ce texte indique, avec un luxe de précisions, comment construire une arbalète montée sur pied, la chirobaliste ; il donne même les dimensions au centimètre près. Le souverain voulut voir ce que valait cet écrit et il fit donc fabriquer une chirobaliste. L'engin achevé, il ordonna à l'un de ses soldats de l'utiliser. L'expérience eut lieu en 1902 à la Saalburg, sur l'emplacement d'un ancien camp romain. Une première flèche, tirée de 50 mètres et par temps calme, atteignit le centre de la cible. Une seconde flèche, dans les mêmes conditions, fendit en deux la première...

L'artillerie romaine a de quoi étonner. Balistes, catapultes et onagres projetaient des boulets et des pierres, des flèches et des javalots, voire des poutres, avec une violence inouïe (voir la figure 2). Ces propulseurs étaient capables de décapiter d'un seul jet de pierre trois hommes ou d'en embrocher autant avec une lance. On imagine mal la puissance développée par des engins pourtant simples, puisqu'ils fonctionnaient sur un principe élémentaire : la détorsion d'un écheveau de crins ou de cheveux fortement torsadé à l'aide d'une sorte de treuil à leviers (voir la figure 1). Ils permettaient de déployer une force de 1200 kilogrammes, allant jusqu'à 6000 kilogrammes dans les meilleures conditions, et une vitesse initiale considérable. Avec la pièce la plus grosse disponible, une baliste, on pouvait lancer à 340 mètres une flèche de 60 centimètres ; à cette distance, elle s'enfonçait encore de 30 centimètres dans une pièce de bois.

Nous avons mentionné que l'armement défensif du soldat comprenait casque, cuirasse et bouclier. En ce qui concerne le casque, il est maintenant bien établi que le fer du casque ne doit pas être en contact direct avec la tête du

fantassin sous peine de lui endommager le cuir chevelu, la souffrance étant aggravée en cas de choc. Car un casque a pour but de protéger contre les coups d'épée qui sont abattus sur le crâne du légionnaire, et aussi contre tout ce qui tombe du ciel – flèches, javalots, pierres. Au début de l'époque impériale, le casque était constitué par une demi-calotte de métal. Par la suite, et pour atténuer le choc des projectiles lancés en tir courbe, un protège nuque a été inventé. Par ailleurs, les Romains avaient découvert un effet d'optique : à une certaine distance, des soldats munis de casques à aigrette paraissent plus grands qu'ils ne sont, ce qui décourage les ennemis. D'où les casques à panache.

Chaque soldat, une forteresse

Pour se protéger le torse, les légionnaires avaient recours à une cuirasse. Dans ce cas également, les reconstitutions de l'*Ermine Street Guard* ont montré qu'une solide protection de la peau contre le métal était indispensable. Au début de l'époque impériale, les soldats se contentaient d'une cotte de mailles ; il s'agit d'une sorte de chemise qui tombe jusqu'à mi-cuisses et qui est fabriquée en fil de fer ; c'est, en quelque sorte, un grand pull-over où le fer remplace la laine. La cotte de mailles protège bien le torse contre les coups d'épée portés de taille (si le coup est très violent, le récipiendaire n'aura qu'un hématome). En revanche, l'efficacité est moindre contre les pointes d'épées, contre les flèches et les javalots, bien que l'impact puisse être fortement atténué si le projectile heurte une des plaquettes reliant les fils métalliques de la cotte.

La cotte de mailles n'ayant pas protégé les légionnaires des flèches des Parthes, il fallut trouver une autre parade. Et ce fut la cuirasse articulée, bien connue grâce à la colonne Trajane, monument encore en place au centre de Rome, et grâce à de nombreuses découvertes (voir la figure 4). Elle est formée de plusieurs bandes de métal reliées par des attaches et qui se recouvraient plus ou moins suivant la



5. Le casque des légionnaires était muni d'un protège nuque destiné à protéger le soldat de projectiles tombant d'en haut. Sa forme était aussi étudiée pour permettre au combattant d'avancer au combat incliné vers l'avant, ce qui lui permettait, protégé par son bouclier, d'éviter les traits arrivant de face. Bastion d'une forte-



resse collective, le légionnaire était aussi une petite forteresse individuelle dont le point faible était les jambes. Elles aussi devaient être renforcées par des jambières portées sur les tibias, sous la tunique recouverte de nombreuses lanières de cuir pour protéger sans entraver les mouvements.

posture du soldat. Les liens donnaient une grande souplesse à la cuirasse et les plaques assuraient une bonne protection contre les armes de jet. Contre des ennemis moins bien équipés que les Parthes, on pouvait utiliser la cuirasse à écailles : des lamelles de métal étaient cousues sur une veste faite en matériau plus souple. Il ne faut pas oublier deux impératifs à propos de la cuirasse : elle ne devait pas être trop lourde, pour ne pas handicaper le combattant, et elle ne devait pas coûter trop cher (en effet, dans l'armée romaine, c'était le soldat qui payait ses armes, l'État ne lui donnant rien d'autre que sa solde).

Dernière pièce de l'armement défensif, et entrant en partie dans l'armement offensif, le bouclier avait la forme d'une grande tuile enrobant le corps. Attaché au bras gauche, il portait en son centre une demi-sphère de métal, l'umbo, qui le protégeait contre les flèches et les javelots, tout en lui permettant de porter des coups à l'adversaire au moment du premier contact. Un renforcement de fer, au sommet, permettait aussi de frapper l'ennemi sous le menton.

Au chapitre bien fourni de l'armement offensif, il faut d'abord mentionner les instruments utilisés par les fantassins légers. Les glands de plomb, lancés par des frondes, pouvaient toucher un homme à plusieurs dizaines de mètres et lui causer de graves blessures, surtout s'ils atteignaient le visage. Les flèches étaient travaillées avec soin : leur section pouvait être carrée, triangulaire ou circulaire, suivant la force de pénétration souhaitée. Les ailes qui flanquaient les pointes des flèches étaient soit lancéolées (en forme de pointe de lance), soit à barbules. Dans le premier cas, un homme blessé pouvait extraire la pointe de la chair sans trop de difficultés ; dans le second cas, il fallait une opération chirurgicale, évidemment pratiquée sans anesthésie. De toute façon, les blessés succombaient souvent en quelques jours, atteints par le tétanos ou par une septicémie.

Le glaive, arme redoutable

Les glaives des légionnaires, les *gladii*, étaient terriblement meurtriers. Les raisons exactes de cette efficacité ne sont pas encore bien connues ; sans doute faut-il chercher du côté des artisans qui les fabriquaient. Très effilée, proche de la rapière, leur lame de 70 centimètres permettait de frapper aussi bien de la pointe que du tranchant. Les premiers phalangites macédoniens qui se sont mesurés aux Romains ont été effrayés par l'efficacité de leurs glaives : ils permettaient de détacher un bras du corps, d'ouvrir largement une gorge ou un ventre. Il ne semble pas que le poignard, porté par chaque soldat, ait eu une fonction dans les duels ; il était plutôt utilisé pour achever les blessés.

Quant aux *pila*, les javelots, ils commencent à être connus précisément grâce à l'archéologie expérimentale. Ils étaient constitués par une longue pointe en fer très mince, fichée dans un manche en bois assez court. Utilisé comme arme de jet, le *pilum* se fichait dans un bouclier, ou dans la terre, ou encore dans le corps d'un homme. Dans le premier cas, la pointe se pliait, alourdissant le bouclier, ce qui forçait l'ennemi à se découvrir. Elle se pliait également dans le deuxième cas, et il était impossible de renvoyer le javelot à l'expéditeur, ce qui était pourtant la tradition des combats de l'An-



6. Les javelots, ou pila (*pilum* au singulier), servaient aux légionnaires à forcer les murailles de la forteresse humaine adverse. Ils étaient lancés à courte distance ou plantés à la main. Leur longueur était calculée pour traverser le bouclier et atteindre le corps de l'adversaire. Comme sa pointe se déformait facilement, un *pilum* restait fiché dans le bouclier atteint. S'il s'y prenait habilement, un légionnaire pouvait alors, en marchant sur le manche du *pilum*, priver son adversaire de son bouclier assez longtemps pour l'attaquer. En bas, l'une des lances légères dont se servaient aussi les légionnaires.

tiquité. Dans le troisième cas, elle causait des blessures profondes. Ces armes pouvaient en outre être utilisées dans une sorte d'escrime qui précédait de peu le corps à corps proprement dit, si le soldat conservait au moins un *pilum* pour cet avant-dernier contact avec l'ennemi.

On sait donc bien que le couple *gladius-pilum* a joué un rôle déterminant dans les succès remportés par l'armée des Romains. Mais des travaux récents, surtout ceux de Guillaume Renoux, qui a soutenu en décembre 2006, à l'Université de Toulouse-le-Mirail, une thèse sur les flèches romaines, ont montré que, grâce à leurs artisans (qui suivaient les légions), les légionnaires ont su se doter d'un armement très efficace, qui les protégeait le mieux possible et qui leur permettait d'infliger de graves blessures. Certes, le champ des recherches en ce domaine n'est pas totalement défriché ; il reste encore des enquêtes à faire. Et surtout, nous ignorons pourquoi, au cours des guerres du III^e siècle, le *gladius* et le *pilum* ont été abandonnés au profit de la longue épée nommée *spatha* et de la lance. Sans doute fallait-il s'adapter à de nouveaux ennemis.

Yann LE BOHEC, professeur à l'Université Paris IV-Sorbonne, est un spécialiste de l'histoire romaine.

M. C. BISHOP et J. C. N. COULSTON, *Roman military equipment from the Punic wars to the fall of Rome*, 2^e édition, Oxbow Books, 2005.

G. BRIZZI, *Le guerrier de l'Antiquité classique, de l'hoplite au légionnaire*, Éditions du Rocher, 2004.

Y. LE BOHEC, *L'armée romaine sous le Haut-Empire*, Picard, 2002.

M. FEUGÈRE, *Les armes des Romains*, Errance, 1993.

M. REDDÉ, *Mare nostrum*, BEFAR, vol. 260, 1986.

<http://www.romans-in-britain.org.uk>

<http://www.vetonia.de/>

Ginkgo biloba : le rescapé et son algue

*L'arbre a traversé les millénaires presque inchangé.
C'est un organisme d'exception, tant par son mode de reproduction
que par l'algue que ses cellules hébergent.*

Jocelyne Trémouillaux-Guiller



*La feuille de cet arbre,
Qu'à mon jardin confia l'Orient
Laisse entrevoir son sens secret
Au sage qui sait s'en saisir.*

*Serait-ce là un être unique
Qui de lui-même s'est déchiré ?
Ou bien deux qui se sont choisis
Et qui ne veulent être qu'un ?*

*Répondant à cette question
J'ai percé le sens de l'énigme
Ne sens-tu pas d'après mon chant
Que je suis un et pourtant deux ?*
Goethe (1749-1832), *Ginkgo biloba*

À la fin de la Seconde Guerre mondiale, le 6 août 1945, une bombe atomique est lâchée sur la ville d'Hiroshima au Japon. À un kilomètre de l'épicentre, un vieil arbre se dresse près du temple d'Housenbou. Le temple est détruit, l'arbre est calciné, tout est mort. Un an plus tard, aucune vie ne reprend sur cette terre irradiée, hormis une petite pousse qui sort du sol à partir de la souche de l'arbre. De cette petite branche, un arbre renaît de ses cendres sans malformation apparente. Après la guerre, la reconstruction du temple nécessitait la transplantation de la souche âgée de 150 ans, un transfert qui mettait en péril l'unique survivant du désastre nucléaire. L'architecte trouva une parade en dessinant les marches de l'escalier du temple en forme de U, ménageant ainsi un emplacement pour l'arbre. Le rescapé est un *Ginkgo biloba*, aussi nommé « arbre de vie » au Tibet ou encore « arbre aux quarante écus » en France.

Nous verrons qu'il a survécu à bien des vicissitudes, car *Ginkgo biloba* L. (le L. signifie qu'il doit son nom d'es-

pèce – *biloba* – à Linné) est le seul du genre *Ginkgo*, de la famille des Ginkgoaceae, de l'ordre ancestral des Ginkgoales et de la classe des *Ginkgoopsida*, un des groupes majeurs des *Ginkgophyta*. L'arbre le plus ancien du monde vivant aujourd'hui représente à lui seul une espèce, un genre, une famille, un ordre et une classe botanique entière. Ce n'est pas sa seule particularité : nous détaillerons comment son mode de reproduction, ainsi que l'algue hébergée par ses cellules, en font un organisme d'exception.

À travers les millénaires

Des fossiles de feuilles et d'autres organes de membres de la classe des *Ginkgoopsida* ont été retrouvés dans des roches du Carbonifère récent d'environ 300 millions d'années. Durant la période précédente, de 410 à 360 millions d'années, les Préspérmatophytes (*Ginkgoopsida*, *Cycadopsida*...) et les Progymnospermes, des plantes qui fabriquent des ovules, apparaissent sur les terres émergées préalablement colonisées par des bryophytes (des mousses) et des ptéridophytes (des fougères arborescentes). Cette végétation est parcourue par des reptiles mammaliens et des amphibiens géants. À la fin de l'ère primaire (de 545 à 245 millions d'années), les marais et les lacs s'assèchent peu à peu, puis au Permien (de 286 à 245 millions d'années), les terres émergées se soudent en un unique continent, la Pangée. Des éruptions volcaniques obscurcissent le ciel et arrêtent les rayons du soleil, déclenchant une épouvantable période glaciaire pendant laquelle se forment une épaisse banquise et des glaciers. Les perpétuels effondrements de ces derniers donnent naissance à de volumineux icebergs, lesquels, poussés par les courants marins, refroidissent les eaux jusqu'aux tropiques. De telles conditions climatiques entraînent la disparition massive d'organismes terrestres et marins : parmi les survivants, figurent les ancêtres des dinosaures et des mammifères.



© Shutterstock/C. Erpenbeck

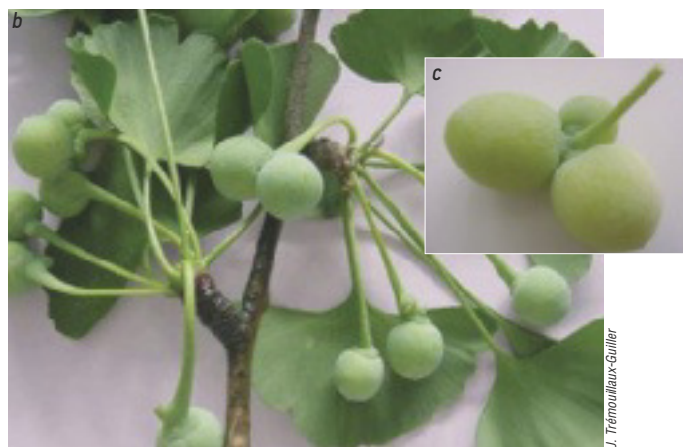


J.-J. Milan

1. Le *Ginkgo biloba* est un arbre dont les ancêtres sont apparus il y a quelque 300 millions d'années, avant la fin de l'ère primaire. Bien que son port (*ci-contre*) ressemble à celui des arbres modernes, il a conservé des caractéristiques ancestrales, comme en témoignent la forme de ses feuilles bilobées et celle des nervures dites dichotomiques.



Steve Crook



2. Le ginkgo est dioïque : l'arbre est mâle ou femelle. Chez le premier, au printemps, des chatons portent des étamines (a) qui libèrent les grains de pollen. Transportés par le vent, ceux-ci sont recueillis sur

les ovules matures, souvent seuls au sommet d'un pédoncule, exceptionnellement groupés par deux ou trois (b et c). La fécondation a lieu quand l'ovule est mature, plusieurs mois après la pollinisation.

Après cette grande crise écologique, dix millions d'années sont nécessaires à la reformation des écosystèmes. Les Gymnospermes, des plantes, tels les conifères, à ovule nu, qui ont remplacé les spores par des graines, ont survécu aux extinctions du Permien et s'épanouissent du Trias (245 à 205 millions d'années) au Crétacé inférieur (il y a 110 millions d'années). Les Ginkgoales apparaissent à la fin du Permien et au début de l'ère secondaire (elle s'étend de 245 à 65 millions d'années), quand les conditions climatiques clémentes favorisent l'épanouissement d'une abondante végétation.

Un « véritable fossile vivant »

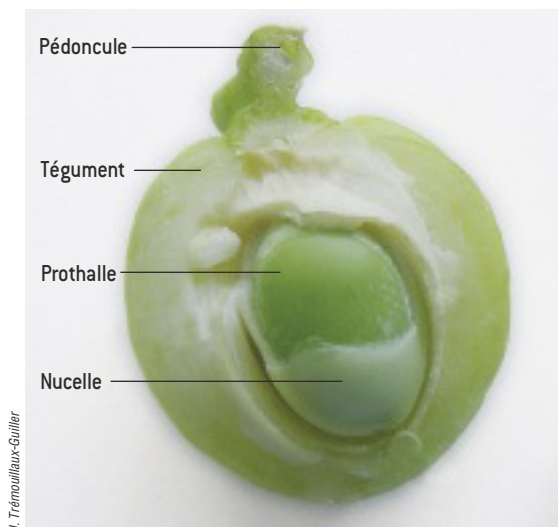
Les Ginkgoales atteignent rapidement leur apogée dans ce paysage peuplé de dinosaures. Des études paléontologiques montrent qu'au Jurassique (205 à 135 millions d'années), le genre *Ginkgo* se diversifie en plusieurs espèces et atteint son pic de diversité au Crétacé (135 à 65 millions d'années) où, dans tout l'hémisphère Nord, de nombreuses espèces sont représentées, dont *Ginkgo adiantoides*. D'après la forme de ses feuilles et la structure de ses ovules, cette espèce serait l'ancêtre de *G. biloba*. Récemment découverts dans la province du Henan, en Chine, des fossiles de 180 millions d'années seraient les plus précoces représentants du genre *Ginkgo*. Cette nouvelle espèce, nommée *G. yimaensis*, diffère de *G. biloba* par des feuilles plus découpées et des ovules plus petits.

Vers 110 millions d'années, une compétition serrée se profile entre les Ginkgoales, les Gymnospermes et de nouveaux venus, les Angiospermes, ou plantes à fleurs, dont les ovules sont protégés par un fruit. De la fin du Jurassique au Crétacé inférieur, des changements géologiques et climatiques bouleversent encore notre planète. La Pangée se disloque et une météorite géante plonge de nouveau la Terre dans l'obscurité et dans un très long hiver. Les dinosaures et d'autres espèces – animales et végétales – disparaissent. Les fossiles révèlent que la diversité et la distribution du genre *Ginkgo* diminuaient déjà. À la fin du Crétacé et au début de l'ère tertiaire (65 à 50 millions d'années), le déclin des ginkgos s'accélère sous l'effet des glaciations successives et de la rivalité des Angiospermes qui dominent le monde végétal.

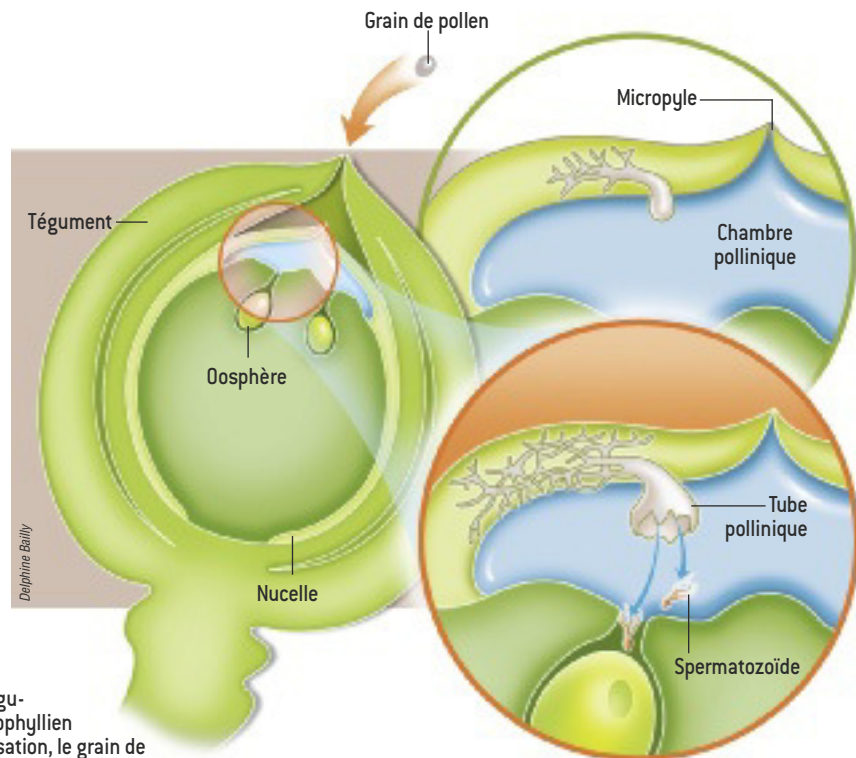
Les grandes glaciations de la fin du Tertiaire et du Quaternaire sont fatales aux ginkgos. La banquise avance sur toute l'Europe et entraîne la disparition des ginkgos il y a 2,5 millions d'années sur ce continent, les chaînes de montagnes orientées d'Ouest en Est constituant un barrage infranchissable vers le Sud. Une seule espèce en réchappe, *Ginkgo biloba*, l'arbre que Darwin qualifiait de « véritable fossile vivant ». L'unique survivant du genre trouve refuge dans les vallées profondes, chaudes et humides, du Sud-Est de la Chine. On le retrouve bien plus tard dans cette région, dans les jardins des monastères et des palais où, pendant plusieurs siècles avant notre ère, il a offert aux empereurs et aux moines bouddhistes sa beauté et ses vertus médicinales (nous y reviendrons). De nos jours, d'anciens spécimens croissent dans des régions centrales de la Chine, notamment la province de Guizhou. On ignore si *G. biloba* existe encore à l'état sauvage, bien que des témoins en aient repéré dans les provinces chinoises de Zhejiang et de l'Anhui.

Au XII^e siècle, *G. biloba* est introduit au Japon et en Corée par des moines bouddhistes revenant de Chine. La quasi-totalité des ginkgos modernes présents en Europe et en Amérique du Nord sont les descendants de quelques arbres japonais. En 1691, le botaniste et médecin allemand Engelbert Kaempfer décrit dans son *Amoenitatum exoticarum* l'arbre qu'il a découvert au Japon. De retour en Europe en 1712, il rapporte dans ses bagages *Ginkyo*, de son ancien nom japonais qui deviendra *Ginkgo* à la suite d'une erreur de transcription transmise à la postérité par Linné. Le naturaliste suédois le nomme *Ginkgo biloba* en raison de la forme bilobée de ses feuilles. D'un vert brillant virant au jaune d'or en automne, sa feuille échancrée ou en forme d'éventail est dotée de nervures dichotomiques, un caractère ancestral (voir la figure 1). Les feuilles fossilisées du Tertiaire ressemblent à celles de nos jardins.

De son lointain passé, *G. biloba* a conservé des caractéristiques botaniques acquises au Jurassique. Certains botanistes le classent d'ailleurs parmi les Préspermaphytes, marquant ainsi la transition entre les Ptéridophytes et les Spermaphytes (les plantes à graines, c'est-à-dire les gymnospermes et les Angiospermes). D'autres relient le ginkgo



J. Trémoullaux-Guiller



Delphine Bailly

3. L'ovule mature de *Ginkgo biloba* est constitué d'un tégument (une enveloppe externe), du nucelle et du prothalle chlorophyllien qui abrite l'oosphère, le gamète femelle. Au moment de la pollinisation, le grain de pollen tombe sur la surface de l'ovule et pénètre par le micropyle dans la chambre pollinique, une cavité emplie d'eau sécrétée par les cellules de l'ovule. Là, il se développe et libère, lors de la fécondation, deux spermatozoïdes, multiflagellés, dans la chambre pollinique. L'un des deux spermatozoïdes (le plus rapide) nage vers l'oosphère avec laquelle il fusionne en une cellule œuf, un futur arbre.

aux gymnospermes, bien que des analyses du génome établissent une parenté plus proche avec les cycadales (les cycas). De fait, ginkgo brouille les pistes : si son port ainsi que ses feuilles bien formées et caduques sont typiques des feuillus, il garde de ses ancêtres une fécondation inféodée au milieu liquide. Voyons comment *G. biloba* s'y prend pour assurer sa descendance.

Le ginkgo est dioïque, c'est-à-dire qu'il existe des arbres mâles et d'autres femelles. Les individus étant souvent éloignés les uns des autres, ils assurent leur descendance grâce à un grain de pollen disséminé par le vent, transportant à distance les gamètes mâles (les cellules reproductrices) vers les organes femelles. *G. biloba* est un arbre ramifié doté de branches dimorphes avec des rameaux longs (des auxilablastes) et des rameaux courts (des mésoblastes), ces derniers portant les organes reproducteurs.

« L'arbre qui pond des œufs »

Au printemps, les inflorescences mâles apparaissent. Ces chatons cylindriques et jaunes, qui ressemblent à ceux des conifères, sont constitués de nombreuses étamines d'où s'échappent des grains de pollen au moment de la pollinisation (voir la figure 2). Chaque grain, d'une forme ovoïde, renferme quatre cellules protégées par deux parois.

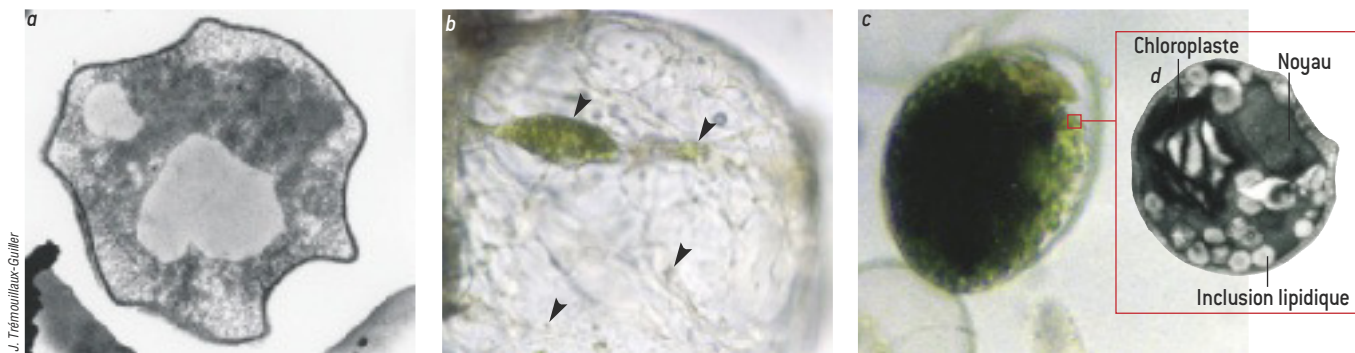
L'inflorescence femelle est composée d'un pédoncule qui porte deux petits ovules droits et nus, l'un des deux avortant le plus souvent. Cependant, on observe parfois des ovules matures groupés par trois sur un même pédoncule. Dans les ovules jeunes, un tégument enveloppe un tissu, nommé nucelle, dont les cellules sont diploïdes, c'est-à-dire qu'elles ont $2n$ chromosomes (ici, $n = 12$). Dans ce nucelle, une cellule donne naissance, après la méiose (la réduction du nombre de chromosomes), à un groupe nommé tétrade de quatre cellules haploïdes (pourvues de n chromosomes).

Puis une cellule de la tétrade se divise et conduit au prothalle femelle, un tissu haploïde où se creusent deux cavités. Dans chacune d'elles se développe une oosphère, c'est-à-dire un gamète femelle ; un seul sera généralement fécondé. À maturité, l'oosphère est une cellule de 0,5 millimètre de diamètre à n chromosomes.

Au moment de sa formation, d'avril à début mai, le prothalle est cœnocytaire : il ressemble à une cellule géante contenant plusieurs noyaux issus de divisions nucléaires successives sans division cellulaire. Ce stade polynucléé s'achève lorsque le nombre de noyaux atteint 256 ! Le prothalle, protecteur et riche en substances de réserves, demeure chlorophyllien bien qu'il soit totalement privé de lumière. On ignore pourquoi. À la fin de l'été, l'ovule mature de *G. biloba* est constitué, de l'extérieur vers l'intérieur, des trois couches distinctes du tégument, du nucelle et du prothalle femelle au sommet duquel se trouvent les oosphères (voir la figure 3).

Lors de la pollinisation, le grain de pollen tombe sur l'ovule au niveau d'une minuscule entrée, le micropyle, et y est piégé par une gouttelette mucilagineuse. Il pénètre ensuite dans la chambre pollinique (une cavité du nucelle) où il patiente quatre mois pendant lesquels il se divise et libère en milieu aqueux deux spermatozoïdes. Ces gamètes mâles (à n chromosomes) sont munis de cils vibratiles et traversent la chambre pollinique. Enfin, un seul spermatozoïde pénètre dans un archégone et fusionne avec l'oosphère pour former la cellule œuf (ou zygote) à $2n$ chromosomes.

Au début de l'hiver, le zygote évolue en un pro-embryon, une cellule d'abord cœnocytaire, puisqu'elle contient elle aussi jusqu'à 256 noyaux avant de se cloisonner. Dans l'ovule mûr, l'embryon, protégé par le prothalle et par le tégument, se développe tout de suite sans nécessiter de phase de dormance, à l'inverse des graines des Spermatophytes. Après une longue adolescence, *G. biloba* n'atteint sa maturité sexuelle que vers l'âge de 20 à 30 ans.



4. Une algue endosymbiotique est hébergée dans des cellules de *G. biloba* sous la forme d'un « précurseur » immature (a). On distingue notamment dans le cytoplasme homogène des masses sombres qui constitueront les membranes du futur chloroplaste. Le précurseur fait trois à cinq micromètres de diamètre. Lorsque la cellule de ginkgo se nécrose, on observe au microscope photonique des algues (b, flèches noires) qui se

divisent en différents points du cytoplasme jusqu'à y proliférer (c). Aucune algue n'est visible dans le milieu extracellulaire, avant l'éclatement des cellules hôtes dû à la multitude d'algues. À maturité, l'algue est dotée de tous les éléments d'une cellule végétale eucaryote (d) : un noyau ici excentré, un chloroplaste et de nombreuses petites inclusions lipidiques. L'algue mesure quatre micromètres de diamètre.

En 1896, Sakugoro Hirase, de l'Université impériale de Tokyo, découvre la mobilité du spermatozoïde multiflagellé de *G. biloba* dans le liquide sécrété par l'ovule. Ce faisant, il démontre la position intermédiaire de cet arbre entre les fougères et les conifères. Chez les premières, la fécondation aquatique se déroule à l'extérieur de la plante, alors que chez le ginkgo la fécondation est interne.

À cause de son mode de reproduction dit ovipare, certains botanistes rapprochent le ginkgo du règne animal. En effet, ses ovules, comme ceux des oiseaux et des reptiles, accumulent des réserves nutritives en l'absence de toute fécondation (chez les plantes à fleurs, l'ovule ne grossit qu'après la fécondation).

L'intérieur des cellules du ginkgo est aussi riche en surprises que l'est son mode de reproduction. D'abord, les nucléoles (des structures sombres du noyau où sont fabriqués les éléments du ribosome) de *G. biloba* sont réticulés comme ceux des cellules animales. Ensuite, et c'est le plus étonnant, ce colosse au port majestueux, de 25 à 40 mètres de hauteur, héberge à l'intérieur de ses cellules une petite algue verte unicellulaire et eucaryote. Une telle association, rare, entre un micro-organisme internalisé et son hôte, est une endosymbiose. Voyons comment l'algue a été découverte.

Une algue masquée

En 1992, j'observais, au microscope, à l'intérieur de cellules âgées en cours de nécrose, des éléments ressemblant à des amas de chloroplastes, ces organites où a lieu la photosynthèse. Plus tard, je notais que, en culture, des cellules de ginkgo privées de paroi (par traitement enzymatique) se nécrosent en quelques semaines. À leur place, d'étranges formations sphériques d'un vert brillant apparaissent. La modification des conditions du milieu entraîna la libération d'algues unicellulaires capables de croître *in vitro*. Cependant, ce premier échantillon ne suffisait pas à prouver qu'une algue intracellulaire de ginkgo était isolée. Il fallait éliminer la possibilité d'une contamination exogène.

L'étude en microscopie électronique à transmission (à l'Université de Tours, au Département de biologie cellulaire

de la Faculté de médecine) de l'algue de *Ginkgo* (de trois à quatre micromètres de diamètre) permettait de lui conférer son statut d'eucaryote. En effet, cette algue a un noyau, des mitochondries et un volumineux chloroplaste, des organites absents chez les procaryotes.

Enfin, en 1998, la division d'algues à l'intérieur de cellules de ginkgo se nécrosant confirmait l'origine endogène de l'algue. Pendant cette intense prolifération intracellulaire, aucune algue ne pouvait être détectée dans le milieu extracellulaire, excluant toute possibilité de contamination externe. En 2000, je réussissais à observer le développement précoce d'algues à l'intérieur d'un cytoplasme encore intact de cellules de ginkgo.

À la même date, j'engageais une collaboration avec Volker Huss, de l'Université d'Erlangen, en Allemagne. Il identifia l'algue en analysant des régions de l'ADN qui codent la petite sous-unité des ribosomes : c'est une *Coccomyxa*, une espèce répandue dans les lichens, des végétaux issus de symbioses mutualistes entre un champignon et une algue unicellulaire. *Coccomyxa* est aussi un parasite de mollusques bivalves. L'association de *G. biloba* et de *Coccomyxa* est un nouveau type de symbiose intracellulaire héritée verticalement, c'est-à-dire de génération en génération, entre deux partenaires eucaryotes. Mieux, les symbioses entre des algues microscopiques et soit des organismes unicellulaires, soit des animaux invertébrés sont répandues, mais elles étaient inconnues chez des végétaux supérieurs. Le ginkgo marque encore sa différence.

L'algue réside à l'intérieur de cellules de ginkgo sous des formes transitoires immatures, dites précurseurs : aucun organe n'est visible dans leur cytoplasme homogène. Seules des masses compactes de nature protéique, plus ou moins dispersées, sont accolées à des inclusions lipidiques. Progressivement, ces masses s'organisent pour former des thylakoïdes, les membranes internes du futur chloroplaste de l'algue. Ce scénario d'édification membranaire est surprenant, car les plastides, de par leur origine également endosymbiotique, ont un mode de division emprunté aux bactéries : leur formation résulte de la division de plastides préexistants.

Quel avantage procure l'algue, dépourvue d'activité photosynthétique dans les cellules chlorophylliennes de

ginkgo ? Les inclusions lipidiques ou lipoprotéiques des précurseurs et des cellules de ginkgo indiquent un rôle possible de *Coccomyxa* dans le métabolisme lipidique de l'hôte. De son côté l'algue, incapable d'assimiler les nitrates, assimile l'azote sous une forme réduite sans doute « empruntée » à son hôte durant sa vie cryptique.

La vie cryptique de l'algue s'achève dans une cellule âgée qui se nécrose. Les répressions exercées par le génome de l'hôte sur cet intrus toléré sont probablement abolies lors de son état prénécrotique, ce qui permet à l'algue de recouvrer son autonomie et de proliférer avant son éjection à l'extérieur. De même, l'algue d'un lichen peut abandonner son partenaire quand leur association cesse d'être avantageuse. Autre cas, les zooxanthelles, des algues symbiotiques, sont expulsées des coraux qui les hébergent quand ils meurent. La libération d'algues de ginkgo requiert sans doute des conditions physico-chimiques particulières, encore inconnues, liées par exemple au vieillissement de l'hôte.

Nos récentes analyses moléculaires ont révélé que l'algue n'est pas uniquement présente dans les arbres du Jardin botanique de Tours qui ont fourni les cellules étudiées. Nous avons recherché des fragments d'ADN codant la petite sous-unité des ribosomes chez l'algue dans différents tissus de ginkgos géographiquement éloignés. Les fragments d'ADN ont été retrouvés dans la quasi-totalité des tissus (reproducteurs, embryonnaires et végétatifs) d'arbres provenant de France, d'Allemagne, d'Italie, du Japon, de Chine et des États-Unis. Ces fragments sont tous identiques à celui de l'algue.

Un ancêtre commun unique

Les ginkgos modernes ont une histoire évolutive d'environ 300 ans, puisqu'ils dérivent tous d'arbres japonais importés. Nous avons analysé des séquences d'ADN nommées ITS1 de l'ADN ribosomal de précurseurs de l'algue présents dans les tissus d'arbres géographiquement éloignés. Ces séquences – des espaceurs intragéniques transcrits, ou ITS – séparent les différentes régions de l'ADN codant les constituants des ribosomes et sont très variables : elles sont utilisées pour reconstituer la phylogénie d'individus étroitement apparentés. Tous les ITS1 étudiés sont identiques à l'ITS1 de l'algue cultivée, et l'on en déduit que les endosymbiontes des *G. biloba* analysés dérivent d'un ancêtre commun.

Comment l'algue est-elle devenue une partie intégrante de *G. biloba* ? Une hypothèse est fondée sur le processus unique de fécondation de cet arbre. Par hasard, une *Coccomyxa* serait tombée sur la gouttelette mucilagineuse qui perle du micropyle et aurait accompagné le pollen. L'algue, comme le pollen, aurait alors patienté quatre mois dans la chambre pollinique où elle aurait trouvé « gîte et couverts » sous forme de cellules nécrosées de l'ovule. Lors de la fécondation, le spermatozoïde (20 fois plus gros que l'algue), une fois expulsé du pollen, aurait entraîné *Coccomyxa* dans l'oosphère. Internalisée dans la cellule œuf, l'algue a pu échapper à la digestion et ajuster sa division afin de continuer sa vie cryptique dans l'hôte. Ce processus ayant conduit à une endosymbiose stable est un événement rare, mais que l'on trouve sporadiquement dans la nature. Ainsi, les mitochondries

que nos cellules hébergent résultent d'une lointaine endosymbiose bactérienne.

Certains prêtent au ginkgo des vertus médicinales. Qu'en est-il vraiment ? Dès 2737 avant notre ère, la médecine traditionnelle chinoise rapportait les bienfaits de cet arbre. L'empereur Chen Nong, dans un recueil sur les plantes médicinales, répertoriait ginkgo comme un remède efficace pour stimuler la circulation sanguine et lutter contre la toux. Plus tard, en 1436 ans avant notre ère (durant le règne de la dynastie Ming), son usage était mentionné pour le traitement de désordres inflammatoires. Introduit au Japon par les moines zen, le ginkgo a été rapidement intégré à la pharmacopée locale.

Une plante médicinale ?

La feuille de *G. biloba* contient de nombreux flavonoïdes et terpènes, dont plusieurs sont spécifiques de l'espèce. En 1932, le chimiste japonais S. Furakawa isolait des feuilles les ginkgolides, des terpènes trilactones uniques dans le règne végétal. Trois décennies plus tard, son compatriote K. Nakanishi identifiait leur structure complexe, et, en 1990, l'Américain Elias Corey, de l'Université Harvard, recevait le prix Nobel de chimie en partie pour la synthèse du ginkgolide B. Ce dernier est un puissant antagoniste du facteur d'activation plaquettaire (une molécule qui déclenche l'agrégation des plaquettes du sang lors de la coagulation du sang) et est efficace pour lutter contre l'inflammation des tissus.

L'extrait standardisé de *G. biloba*, issu des feuilles, contient 24 pour cent de flavonoïdes (des antioxydants qui piègent les radicaux libres délétères) et 6 pour cent de ginkgolides et de bilobalide efficaces pour traiter certaines pathologies liées au vieillissement et des désordres vasculaires. Ces dernières années, notamment en 2007, de nombreux résultats sur les activités pharmacologiques des extraits de ginkgo ont été publiés dans des revues scientifiques internationales.

Le ginkgo a une résistance naturelle aux insectes ravageurs, aux champignons, à la pollution atmosphérique et industrielle ainsi qu'au nucléaire. Pour cette robustesse, il est souvent choisi pour border les avenues, les places et les squares de grandes villes. Ainsi, à New York, tout arbre qui meurt est remplacé par un *G. biloba*. Cette résistance explique sans doute la longévité de cet arbre qui peut atteindre l'âge de 1 500 ans, voire beaucoup plus selon certains. Toutefois, le ginkgo est classé parmi les espèces végétales en péril sur notre planète. Il méritait bien d'être élu, en France, *Arbre de l'An 2000* !

Jocelyne TRÉMOUILLAUX-GUILLER est professeur honoraire de biologie à l'Université François Rabelais, à Tours.

J. TRÉMOUILLAUX-GUILLER et V. A. R. HUSS, *A cryptic intracellular green alga in Ginkgo biloba: ribosomal DNA markers reveal worldwide distribution*, in *Planta*, vol. 226, pp. 553-557, 2007.

J. TRÉMOUILLAUX-GUILLER et al., *Discovery of an endophytic alga in Ginkgo biloba*, in *American Journal of Botany*, vol. 89, n° 5, pp. 727-733, 2002.

T. HORI, R. W. RIDGE, W. TULECKE, P. DEL TREDICI, J. TRÉMOUILLAUX-GUILLER et H. TOBE (Eds.), *Ginkgo biloba – A global treasure, from biology to medicine*, Springer-Verlag Tokyo, 1997.

Témoins d'effraction : les scellés

Utilisés depuis des millénaires, les scellés remplissent de multiples et essentielles fonctions de sécurité. Mais de nouveaux développements devraient améliorer leur efficacité, aujourd'hui assez médiocre.

Roger Johnston

Depuis la nuit des temps, l'homme se pose de grandes questions : « Qui suis-je ? Comment le monde a-t-il été créé ? Quel est le sens de la vie ? » Et depuis tout aussi longtemps, une autre question un peu plus terre à terre le tourmente : « Quelqu'un aurait-il mis son nez dans mes affaires ? »

Les archéologues savent qu'il y a au moins 7000 ans, les hommes cherchaient déjà à assurer la sécurité de leurs biens. Les musées du monde entier contiennent des exemples de dispositifs conçus pour détecter une éventuelle effraction ou intrusion.

Les questions d'accès et de manipulation non autorisés sont aujourd'hui une préoccupation majeure dans tous les domaines de la vie quotidienne. Elles concernent au premier chef les divers aspects de la sécurité publique, comme la recherche d'armes et d'explosifs dans les ports et aéroports, la surveillance des installations nucléaires, le transport et le stockage de matières radioactives et de produits chimiques dangereux, ou encore la sécurisation des urnes et des machines à voter.

L'enjeu est tout aussi important pour les consommateurs, avec la sécurité des aliments, des médicaments ou de l'archivage, numérique ou sur papier. Il s'agit aussi de se prémunir contre l'utilisation ou la manipulation inappropriées d'équipements médicaux, de garantir la bonne calibration d'instruments de mesure, de s'assurer que les compteurs d'eau, d'électricité ou de gaz n'aient pas été trafiqués, que sacs postaux ou pièces à conviction de la justice n'aient pas été ouverts, et ainsi de suite.

Un autre domaine majeur, bien que peu apparent, est le transport des marchandises. Les experts de la sécurité estiment que deux pour cent du fret mondial est volé. Aux États-Unis, ces pertes pourraient s'élever à environ 50 milliards de dollars par an, les estimations allant de 2 à 150 milliards de dollars. Les transporteurs doivent également se méfier des trafiquants de drogue, qui ne volent pas

la marchandise, mais l'utilisent pour cacher et acheminer leurs produits illicites.

Un système de protection complet contre les effractions est composé de nombreux éléments. Je me concentrerai ici sur l'un d'eux : les scellés, c'est-à-dire les dispositifs conçus pour révéler l'existence d'une effraction. Pour comprendre comment ils fonctionnent, il est utile de les comparer avec d'autres types d'équipements de sécurité.

Considérons les serrures ou les cadenas, dont la vocation est en général uniquement de retarder, de compliquer et de décourager les intrusions. Parfois, les serrures s'accompagnent de détecteurs d'intrusion, en d'autres termes d'alarmes anticambriolage. Le plus souvent, ces dispositifs transmettent un signal d'alerte à la police ou à des agents de sécurité pour qu'ils se rendent sur les lieux de l'effraction, afin d'appréhender les intrus ou, au moins, de les mettre en fuite avant qu'ils ne commettent davantage de dégâts. Mais ces détecteurs d'intrusion sont coûteux, complexes et ils se déclenchent parfois intempestivement. En outre, ils impliquent une veille de la police ou des agents de sécurité, appelés à intervenir rapidement, et ces dispositifs peuvent être très difficiles à mettre en œuvre sur des chargements en cours de transport.

Donner l'alarme après coup

Il est parfois irréaliste d'essayer d'empêcher les effractions ou d'intervenir rapidement. Dans bien des situations, il suffit juste d'être averti qu'une effraction il y a eu. C'est la raison d'être des scellés. Contrairement aux serrures et cadenas, les scellés ne sont pas censés résister aux effractions ni les retarder sérieusement. Et contrairement aux détecteurs d'intrusion, ils ne donnent pas l'alarme en temps réel.

Les scellés ont simplement pour mission de trahir les accès non autorisés. Ils se présentent souvent sous la forme d'emballages comportant un témoin d'ouverture, tels ceux que l'on trouve sur les biens de consommation (aliments, médicaments). Ils peuvent aussi être conçus comme des scellés bar-



©The Trustees of the British Museum

1. L'écriture cunéiforme sur cette tablette d'argile vieille de 3 700 ans, découverte à la frontière de la Syrie et de la Turquie, détaille une transaction de propriété. Le contrat était inséré dans une enveloppe d'argile, aujourd'hui brisée, qui a été estampillée des sceaux personnels des témoins de la transaction. De tels procédés ont été utilisés durant des milliers

d'années pour garantir l'authenticité et marquer la propriété. Aujourd'hui, on trouve des systèmes allant des films de plastique sur les produits alimentaires aux boucles de fibre optique sur les conteneurs de matières nucléaires. Mais le but des scellés n'a pas changé : il s'agit d'alerter les utilisateurs de toute intrusion.

rières, qui combinent une serrure et un scellé en un même dispositif ; de tels scellés équipent souvent les portes de camions, de wagons ou de conteneurs, pour protéger leur cargaison.

Il est incorrect de dire des scellés qu'ils sont « inviolables » ou même « résistants aux intrusions ». Leur intérêt vient du fait qu'ils sont au contraire faciles à briser, mais très difficiles à réparer.

Pour certaines applications, les serrures et les alarmes ne sont pas adaptées. De tels dispositifs sont trop lourds et encombrants pour équiper les articles en vente libre ou les colis postaux. En revanche, les scellés peuvent être légers, bon marché et jetables. Ces caractéristiques sont également importantes pour les conteneurs, qui peuvent parcourir le monde des années durant avant de revenir à leurs propriétaires d'origine. Les scellés n'ont pas de codes ou de clés à conserver précieusement, leur fonctionnement ne nécessite en général aucune source d'énergie et on peut les retirer rapidement en cas d'urgence : ce sont là des aspects importants pour l'industrie des transports.

Plus de dix millions de scellés seraient ainsi installés ou inspectés tous les jours rien qu'aux États-Unis. Si l'on inclut les emballages à témoin d'ouverture dans le décompte, on dépasse largement les 200 millions par jour. Les scellés sont ainsi une composante importante de la sécurité moderne au quotidien.

Comment l'histoire des scellés a-t-elle commencé ? Il y a des dizaines de milliers d'années, les gens balayaient

probablement le sol devant leur porte, leurs sanctuaires religieux ou leurs réserves de nourriture et d'armes avant de se rendre ailleurs. À leur retour, ils recherchaient les empreintes de pas dans le sable ou la poussière, afin de repérer les traces éventuelles d'un passage indésirable. Bien sûr, les intrus pouvaient facilement contrer ce système et effacer les traces avec une simple branche d'arbre. Et c'est ainsi qu'un jour, l'idée se serait imposée de marquer ces endroits stratégiques avec des dessins difficiles à copier.

Une pratique millénaire

Les hommes auraient commencé à réaliser des scellés personnalisés au Néolithique. Il s'agissait en général d'un petit disque ou cylindre d'argile, de bois, d'os ou de pierre, gravé d'un motif particulier. On l'utilisait, par exemple, pour marquer de façon distinctive le bouchon en terre d'une jarre. Le marquage était fait sur l'argile encore molle soit en appliquant un sceau en forme de disque, soit en faisant rouler un sceau cylindrique autour du bouchon. Quiconque ouvrait le récipient provoquait la destruction du motif appliqué. Ne disposant pas du sceau original, l'intrus était alors bien en peine de reproduire le motif sur un bouchon de remplacement.

L'estampillage de biens avec des sceaux en forme de disque s'est pratiqué pendant des milliers d'années avant l'invention de l'écriture, vers 3200 avant notre ère. De fait,



2. Les sceaux cylindriques mésopotamiens étaient roulés sur des tablettes d'argile, comme ce sceau fait en roche volcanique et datant de 3000 avant notre ère environ (a). Une chevalière en or égyptienne datant de l'an - 500 environ était appliquée sur de petits morceaux d'argile ou de grosses gouttes de cire pour sceller les documents en papyrus (b). Les puissants et les riches avaient souvent des sceaux en pierre décorative, comme cet exemple en cornaline, qui appartenait à l'intendant en chef de l'Iran au V^e siècle de notre ère (c). Un scellé métallique protégeait le cordon qui fermait le parchemin de la correspondance officielle de l'empereur byzantin Andronic II, au XIV^e siècle (d).

certain spécialistes pensent que les symboles dessinés sur les sceaux ont joué un rôle moteur dans le développement de l'écriture comme de l'arithmétique.

Les sceaux cylindriques ont été inventés en Mésopotamie aux alentours de 3500 avant notre ère. Leur avantage était de faire tenir plus de dessins sur un sceau donné, et le motif pouvait être reproduit à l'infini en faisant tourner le sceau sur l'argile. Les premiers sceaux étaient gravés dans de l'argile, du bois ou de l'os relativement tendres, et la plupart ne sont pas parvenus jusqu'à nous. Il existe cependant de nombreux exemples de très anciens sceaux en pierre. Les plus vieux présentent des motifs géométriques simples, mais les dessins gravés sur les sceaux plus récents sont remarquablement élaborés.

Des sceaux aux scellés

Les sceaux anciens remplissaient essentiellement les mêmes fonctions que leurs équivalents d'aujourd'hui. Ils étaient utilisés pour détecter des effractions portant sur un objet ou un lieu, pour garantir la sécurité des marchandises, pour faciliter les inspections douanières. Ils avaient également des fonctions d'étiquetage : garantir l'authenticité, aider aux inventaires, identifier les propriétaires ou la marque, acheminer de l'information (en particulier quand peu de personnes savaient lire ou écrire). Les sceaux anciens et leurs motifs étaient également utilisés pour la décoration, à des fins religieuses ou magiques, et comme signatures légales. Aujourd'hui encore, dans certaines régions d'Asie, le sceau d'un individu fait davantage autorité au regard de la loi ou de la société que sa signature manuscrite.

Peu après 3000 avant notre ère, les Égyptiens utilisaient des morceaux d'argile, appelés *bullae*, et de la ficelle pour sceller les papyrus officiels. Ils plaçaient également des scellés sur les tombes de leurs défunts. Quand la chambre funéraire était achevée et le corps momifié placé à l'intérieur, les bords de la porte étaient recouverts de boue et de plâtre. Il était encore possible d'ouvrir la porte, mais le scellé était alors détruit. Les archéologues savent ainsi par avance si une tombe a été pillée, juste en vérifiant l'intégrité du scellé.

À mesure que les matériaux se sont perfectionnés, les scellés en ont fait de même. À partir du IV^e siècle de notre ère environ, les scellés en plomb se sont répandus. On noue une ficelle autour de l'objet à sceller, puis on passe les deux extrémités de la ficelle à travers un morceau de plomb mou. À l'aide d'une presse manuelle (même principe que le papier gaufré), on coince une partie des deux extrémités de la ficelle dans le métal malléable, en même temps que s'imprime un motif de chaque côté du morceau de plomb. Il est alors impossible d'accéder au contenu de l'objet sans couper la ficelle ou casser le plomb, donc sans trahir son geste.

À partir du XII^e siècle, les Européens ont commencé à utiliser des scellés de cire. Pour marquer un document ou une enveloppe, l'expéditeur laissait couler dessus de la cire fondue et y apposait son sceau personnel. La cire a été remplacée plus tard par de la gomme-laque ou de la résine.

De nos jours encore, certaines personnes aiment embellir leur correspondance en la scellant avec de la cire. Cela paraît anachronique ? Le programme russe de sûreté nucléaire conti-

nue à utiliser des scellés à la cire, même si cette technique est beaucoup moins utilisée que par le passé. Et beaucoup d'autres pays continuent à utiliser des scellés de plomb pour de nombreuses applications touchant à la sécurité.

La nouvelle avancée majeure a eu lieu dans les années 1880 et 1890, quand un certain nombre d'inventeurs ont conçu et breveté, pour les wagons de chemin de fer, divers types de scellés peu coûteux et jetables. Ces innovations ont permis de protéger les wagons de marchandises sans recourir à de grosses serrures lourdes et onéreuses, et de surcroît longues à retirer une fois le train arrivé à destination. Si le scellé était intact à l'arrivée, il n'était pas nécessaire de faire l'inventaire du contenu ou de vérifier que le chargement n'avait pas été pillé.

Le fabricant de tels scellés ayant connu à cette époque le plus gros succès était Emil Tynden, un Suédois émigré aux États-Unis, fondateur en 1897 de l'*International Seal and Lock Company* à Hastings, dans le Michigan. Son usine fut l'une des premières du monde à être automatisées et à produire en masse, une décennie avant que Henry Ford ne monte ses célèbres chaînes d'assemblage pour automobiles.

Les dispositifs inventés par Tynden et d'autres sont très proches des scellés à boucle métallique encore en usage aujourd'hui. Une des extrémités de la bande de

métal ayant été insérée dans l'autre, le scellé est irréversiblement fermé, de façon que la solution la plus simple pour l'ouvrir consiste à le couper. Chaque scellé est gravé d'un numéro de série unique. Si l'on découvre que le scellé est endommagé ou manquant au moment de l'inspection, ou s'il ne présente pas le bon numéro de série, les inspecteurs savent qu'une manipulation a eu lieu.

Dispositifs passifs ou actifs

D'autres types de scellés modernes sont électroniques, ou actifs, et fonctionnent sur piles. Par exemple, dans un scellé actif à fibre optique, une extrémité de la fibre est glissée dans le loquet du conteneur à surveiller, puis réinsérée dans le scellé. Celui-ci envoie périodiquement une impulsion lumineuse le long de la fibre optique. Si le dispositif détecte qu'une impulsion lumineuse n'a pas achevé son parcours en boucle, c'est que la fibre a été coupée ou déconnectée.

Parmi les scellés passifs modernes, on trouve les assemblages mécaniques irréversibles (tels les scellés à boucle métallique des wagons de chemin de fer), ainsi que des scellés très fragiles qui s'endommagent si quelqu'un essaie de les enlever ou de les ouvrir. Les scellés barrières, à l'inverse, supportent des forces allant jusqu'à plusieurs centaines,

Contrefaire des sceaux d'argile

Pendant des milliers d'années, en Mésopotamie et en Égypte ancienne, des sceaux apposés sur de l'argile ont protégé les propriétés et la correspondance.

Quel était le niveau de sécurité de ces sceaux ? Pas très bon. De fait, il existe des preuves de contrefaçon très ancienne, sous la forme de petits cylindres ornés, datant de 2500 avant notre ère environ. Contrairement aux sceaux cylindriques plus typiques de l'époque, faits de pierre, ceux-ci étaient façonnés dans de l'argile. Les motifs complexes y ont peut-être été transférés en faisant rouler les cylindres d'argile sur des motifs plans imprimés par d'autres sceaux. Si c'est le cas, ces cylindres d'argile ont pu servir de contrefaçons, mais nous n'en avons pas la preuve. On connaît également des exemples de contrefaçons probables, anciennes et modernes, par des artisans qui ont gravé des sceaux neufs ressemblant à s'y méprendre à l'original.

Mes collègues du Laboratoire national de Los Alamos et moi-même avons décidé de voir si nous pouvions imiter les empreintes de sceaux sur de l'argile à l'aide de matériaux facilement disponibles il y a plusieurs milliers d'années. Nous avons mis au point plusieurs méthodes, mais la contrefaçon s'est révélée le moyen le plus simple. Nous avons estampillé plusieurs types d'argile avec des sceaux modernes en laiton. Nous avons graissé les empreintes avec de l'huile d'olive ou de sésame, puis nous en avons fait des moulages à l'aide d'argile ou de cire d'abeille. Une fois durcis, ces moulages pouvaient être utilisés pour créer des empreintes contrefaites.

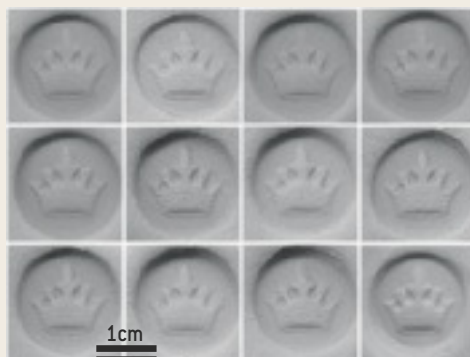
Des volontaires ont inspecté nos œuvres dans le cadre de plusieurs tests. Nous avons montré aux inspecteurs un original, qu'ils ont ensuite pu conserver ou non selon les cas. Nous leur avons alors présenté une seule ou plusieurs

empreintes à la fois et nous leur avons demandé de déterminer lesquelles étaient des faux. En moyenne, les volontaires avaient tort environ une fois sur trois. Et bien qu'on les ait informés que les faux étaient très peu nombreux, ils avaient tendance à en trouver partout.

À leur grande frustration, nous n'avions pas expliqué à nos inspecteurs volontaires comment repérer une contrefaçon. Il est vraisemblable que leurs homologues de l'Antiquité se trouvaient dans la même situation. Mais une comparaison attentive des empreintes contrefaites avec les originaux est un moyen assez fiable de repérer les faux. Les moulages d'argile ou de cire d'abeille que nous avons utilisés pour créer les empreintes contrefaites ont rétréci au séchage. Ainsi, les contrefaçons réalisées à partir de ces moulages étaient plus petites que les estampillages originaux, de 7,5 pour cent en moyenne pour la cire d'abeille, et de 9,5 pour cent avec l'argile, bien que cette diminution modeste puisse passer inaperçue en l'absence de l'original. Sur l'image ci-dessous, l'estampillage contrefait est situé en bas à droite.

Les inspecteurs de sceaux ou scellés modernes sont rarement équipés d'exemplaires de référence, ce qui permet à des faux raisonnablement travaillés de passer inaperçus. Les archéo-

logues ont trouvé de petits rectangles d'argile sur lesquels des sceaux avaient été apposés. Il peut s'agir de jetons de commerce, de cartes d'identité ou de cartes de visite. Mais il pourrait aussi s'agir de références pour identifier ou authentifier un sceau...



Roger Johnston

Ces empreintes sur de l'argile ont été réalisées avec un sceau métallique. Celle du bas à droite résulte d'un sceau contrefait, obtenu grâce à un moulage de cire ou d'argile : le motif imprimé est plus petit.

voire milliers, de kilogrammes. Certains scellés à fibres optiques sont également passifs : une photographie de l'extrémité du faisceau de fibres optiques est comparée à une autre, prise à une date ultérieure. Si le faisceau a été coupé ou remplacé par un autre, les clichés seront différents.

Il est peu fréquent de rencontrer des scellés électroniques au quotidien. Pour la plupart des gens, les emballages à témoin d'ouverture sont la forme la plus courante de scellés. Depuis l'affaire (non résolue) du Tylenol empoisonné, qui a tué huit personnes aux États-Unis en 1982, la loi américaine oblige tous les produits pharmaceutiques en vente libre à être munis d'un témoin d'ouverture approuvé par les autorités (la FDA). De nombreux autres articles industriels et de grande consommation, tels les produits alimentaires, sont équipés de témoins d'ouverture : films plastiques délicats, capuchons et couvercles munis de parties autocassantes, feuille d'aluminium thermocollée ou autre doublure fragile placées sous le couvercle, ou encore « blisters » qui protègent individuellement des comprimés de médicament.

Malheureusement, les témoins d'ouverture actuels ne sont pas très satisfaisants. Les exigences quant à leurs per-

formances et aux tests de ces performances ne sont pas explicitées, et la conception des dispositifs manque d'imagination. Les témoins d'ouverture ne sont pas tant conçus pour détecter efficacement des intrusions que pour se prémunir contre des accusations de négligence et réduire les indemnités de dommages et intérêts en cas de problème.

Par ailleurs, l'étiquetage et les informations concernant les témoins d'ouverture et destinés au consommateur manquent fréquemment de clarté. Les informations sur la protection de l'emballage sont inexistantes ou imprimées en tout petits caractères. Et souvent, la seule indication concernant le scellé est placée sur le scellé lui-même : si le scellé est retiré proprement, il ne reste plus aucun signe qu'il y en avait un.

Silence sur les défauts

La raison d'être de ces défauts est simple : les fabricants ne souhaitent pas vraiment que leurs produits soient associés, dans l'esprit des consommateurs, à l'éventualité de manipulations suspectes. Ils rechignent donc à attirer l'attention sur les témoins d'ouverture de leurs emballages et parlent plutôt de « fermetures fraîcheur ».

Tout cela explique qu'une manipulation frauduleuse ne nécessite, avec un peu de pratique, qu'un outillage assez simple. Une telle manipulation fait généralement intervenir l'un de ces cinq types possibles d'attaque : l'ouverture du récipient ou de l'emballage sans laisser de trace évidente ; l'élimination des témoins d'ouverture en espérant que l'utilisateur final ne s'apercevra pas qu'ils manquent ; l'élimination du témoin d'ouverture et son remplacement par des imitations pour rassurer faussement l'utilisateur plus attentif ; la réparation, l'effaçage ou la dissimulation de toute trace d'intrusion ; la réalisation d'une contrefaçon de l'emballage ou de son témoin d'ouverture.

Les problèmes dont souffrent les témoins d'ouverture ont été étonnamment peu étudiés, même par les fabricants dont la responsabilité légale est fortement engagée si leurs produits subissent des fraudes. Notre équipe a réalisé un petit test lors du septième symposium des scellés de sécurité, tenu en 2006 à Santa Barbara, en Californie. Nous avons présenté aux 72 experts qui assistaient à la conférence un certain nombre de produits alimentaires et pharmaceutiques, et nous leur avons demandé de déterminer lesquels avaient été ouverts. Les effractions avaient été perpétrées rapidement par un de nos étudiants, à l'aide de matériaux et d'outils faciles à se procurer. Les participants à la conférence n'auraient pas obtenu un meilleur score de réussite s'ils avaient désigné les emballages totalement au hasard ! Si la conception des

3. Il existe aujourd'hui plus de 5 000 types de scellés. Certains sont électroniques et utilisent des fibres optiques et des capteurs pour détecter les intrusions (*rangée du haut*). Certains scellés passifs utilisent également des fibres optiques (*les deux boucles noires, deuxième rangée, à droite*). Certains dispositifs, nommés scellés barrières, sont prévus pour servir de cadenas aussi bien que de scellés (*deuxième et troisième rangées, à gauche*). Les scellés à boucle, aujourd'hui essentiellement en plastique, sont utilisés sur les conteneurs de marchandises (*à gauche au milieu*). Les scellés à boucle métallique sont des versions modernes des anciens scellés de plomb (*en bas à gauche*). Les fragiles pellicules de plastique ou de papier sont prévues pour se déchirer ou se déformer irréversiblement à l'ouverture (*au milieu et en bas à droite*).

Roger Johnston



témoins d'ouverture ne permet pas à des professionnels de repérer une intrusion, comment le consommateur ordinaire pourrait-il remarquer un tel acte frauduleux ?

Il existe au moins 105 façons générales différentes de mettre en échec des scellés. Par « mettre en échec », je veux dire retirer le scellé, puis le remettre ou le remplacer par une contrefaçon, sans que cela ne se détecte.

Les attaques se divisent en dix catégories principales. Dans les « attaques par crochetage », un outil spécial est utilisé pour ouvrir le scellé sans l'endommager ou laisser d'indices. Cette approche fonctionne bien sur un nombre surprenant de scellés. Dans les « attaques par déscellement », le dispositif peut être abîmé, mais les traces de l'effraction sont efficacement cachées ou réparées. Si le malfrat a accès au scellé avant son utilisation, il peut utiliser une « attaque par la porte de service » en introduisant un défaut exploitable dans le scellé, au stade de la conception de celui-ci, de sa fabrication, de son acheminement, de son stockage ou juste avant son utilisation. Alternativement, il peut saboter la procédure de scellement, en utilisant quelqu'un de l'intérieur qui posera un scellé inapproprié ou qui ne fermera pas la porte avant le scellement, afin qu'un complice puisse accéder au conteneur, puis le fermer correctement avant son inspection.

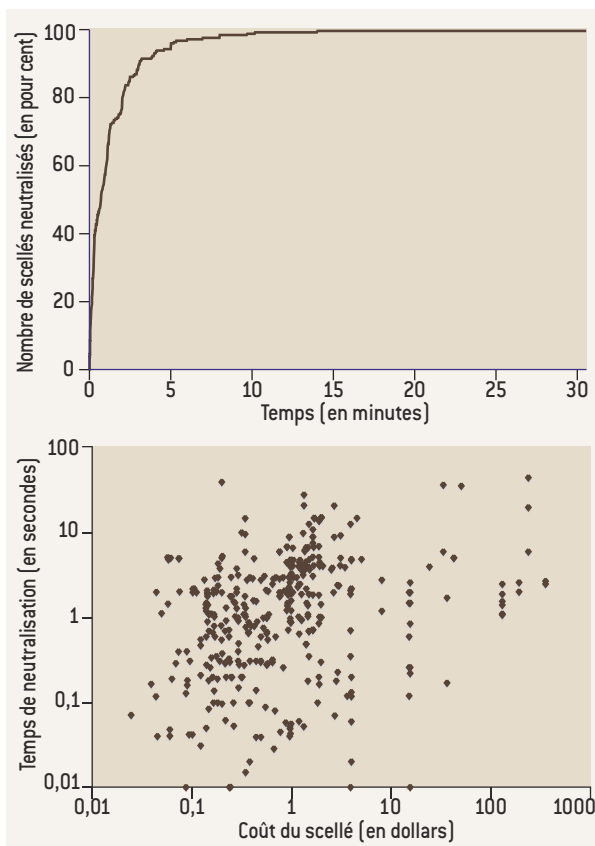
Plus vulnérables qu'on ne le croit

Les transgresseurs peuvent également modifier les données du scellé, tels son numéro de série ou des rapports d'inspection. Si les scellés sont actifs, ils peuvent être victimes (ou leurs lecteurs) d'assauts électroniques sur leurs divers composants (capteurs, logiciels ou données stockées en mémoire). Plus élaboré encore, il peut y avoir des « attaques par duplication », où l'on crée une copie du scellé à l'usine qui les fabrique. Les délinquants doivent alors s'arranger pour que des employés du complexe de production compromettent la sécurité. Il ne faut pas confondre cette stratégie avec la contrefaçon, où un faux (généralement grossier) est fabriqué en dehors de l'usine à partir de nouveaux scellés, de parties usagées ou de matériaux complètement originaux.

Mes collègues de Los Alamos et moi-même avons analysé des centaines de scellés utilisés par le gouvernement ou par les entreprises, depuis les modèles mécaniques les plus élémentaires jusqu'aux dispositifs électroniques de haute technologie.

Nous avons mis en évidence comment ces scellés, dont le coût varie dans une gamme qui va de 1 à 10000, peuvent être rendus inutiles rapidement et facilement grâce à des outils, des fournitures, des méthodes et une dextérité très ordinaires, ressources que tout un chacun peut aisément se procurer. Bien que notre laboratoire nous donne accès aux technologies les plus avancées, nous n'avons encore jamais rencontré un scellé (même parmi ceux utilisés dans le domaine de la sûreté nucléaire) dont la mise en échec nécessite de déployer les grands moyens.

Et ce n'est pas faute d'avoir essayé : nous avons soigneusement étudié 244 types de scellés différents, plus quelques centaines d'autres moins en détail. Parmi les 244 types de scellés étudiés de près, la moitié sont utilisés



4. Des tests effectués sur 244 types de scellés ont montré que la plupart étaient relativement faciles à attaquer. La majorité pouvaient être neutralisés (retirés et remplacés sans laisser de traces) par une personne seule en moins de deux minutes, et tous les dispositifs pouvaient être mis en échec en moins d'une demi-heure (*graphique du haut*). Les scellés les plus onéreux n'étaient pas moins vulnérables : une représentation de 393 attaques différentes sur ces 244 scellés ne fait apparaître qu'une très faible corrélation entre le temps nécessaire à la neutralisation et le coût du scellé (*ci-dessus*).



5. En 1982, aux États-Unis, un médicament (du Tylenol) auquel a été ajouté du cyanure a tué huit personnes. L'affaire n'a jamais été élucidée, mais, depuis, les flacons (*à gauche*) ont été remplacés par d'autres, équipés d'un témoin d'ouverture (*à droite*).

pour ce qui peut être considéré comme des applications critiques, là où un haut niveau de sécurité est requis. Par exemple, 46 d'entre eux étaient utilisés quelque part dans le monde pour le contrôle de la sûreté nucléaire. Nous avons réussi à mettre en échec l'un d'eux de six façons différentes !

Contrairement à ce que l'on pourrait penser, nous avons découvert que les scellés électroniques coûteux ne sont pas beaucoup plus efficaces que les scellés mécaniques bon marché (en tout cas tels qu'ils sont actuellement conçus et utilisés). Les scellés actifs comportent davantage de pièces susceptibles d'être attaquées, et ceux qui les inspectent se fient trop souvent aux seuls lecteurs électroniques : si la machine indique que le scellé est intact, l'inspecteur la croit, même lorsqu'il y a des traces physiques évidentes d'effraction sur le scellé lui-même.

Dans nos tests, la durée moyenne d'une attaque sur un scellé donné était de 1,4 minute, avec une valeur médiane

de 43 secondes seulement. Le coût en était également très bas. Pour la première attaque sur un scellé, les outils et fournitures revenaient en moyenne à 78 dollars (environ 55 euros), la médiane étant à cinq dollars (environ 3,5 euros). Mais une fois l'investissement fait, les attaques suivantes du même type ne coûtaient plus que 62 cents en moyenne chacune, avec une médiane à 5 cents. La statistique la plus éloquente est peut-être qu'il nous fallait seulement 2,3 heures en moyenne pour déterminer la bonne méthode, même s'il fallait souvent beaucoup plus longtemps pour maîtriser la technique. La pratique contribuait également à réduire le temps de planification des attaques.

Lorsque mes collègues et moi-même faisons la démonstration de la vulnérabilité des scellés à des responsables de la sécurité, nombre de ceux-ci se demandent si l'utilisation de ces dispositifs a le moindre intérêt. Nous soulignons alors que ces scellés peuvent être très efficaces, pourvu qu'ils soient utilisés correctement. Pour environ 60 pour cent des attaques de scellés que nous avons élaborées dans notre étude, il existe des contre-mesures simples et peu coûteuses, et pour 30 pour cent des attaques, il existe des parades plus compliquées.

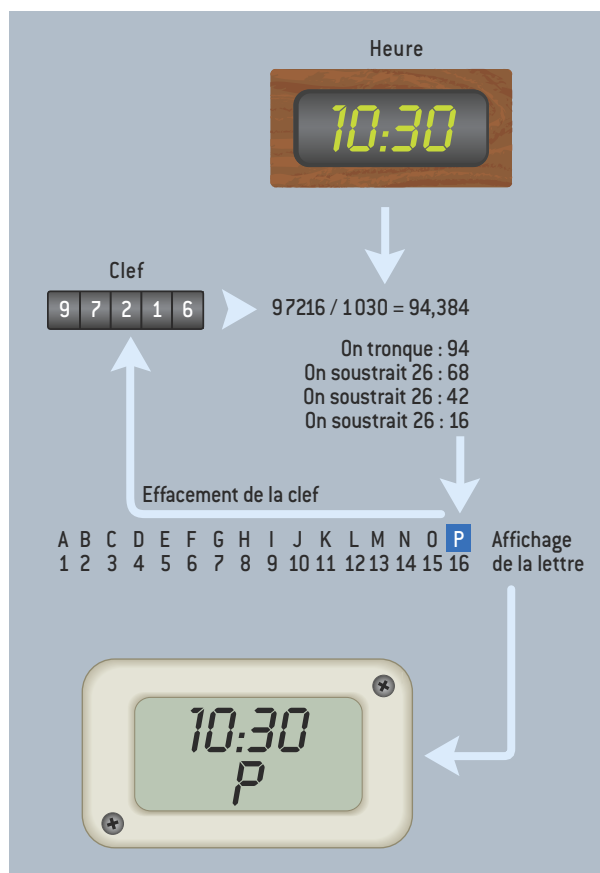
Plus efficaces... si bien utilisés

Il peut s'agir de modifications mineures du scellé, mais la solution consiste le plus souvent à changer les procédures de mise en place et de vérification des scellés. Par exemple, la plupart des inspecteurs ne sont pas munis d'un scellé de référence qui leur permettrait de faire des comparaisons. De nombreuses entreprises ne sont également pas assez prudentes lors de leur mise au rebut de scellés usagés, et mettent involontairement à la disposition des personnes malintentionnées des pièces pouvant servir à la contrefaçon.

Pour que l'utilisation des scellés actuels soit efficace (et cela reste valable pour les scellés électroniques de haute technologie lus avec des lecteurs semi-automatiques), il faudrait que les inspecteurs aient une parfaite compréhension des scellés qu'ils utilisent et de leurs faiblesses potentielles, afin qu'ils puissent rechercher en connaissance de cause les traces d'attaques correspondant aux scénarios les plus vraisemblables.

Une telle expertise nécessite une formation approfondie assortie d'une bonne pratique, des informations détaillées sur les faiblesses des scellés, un solide sens de l'observation et un bon esprit critique. Peu d'utilisateurs de scellés sont prêts à assumer ces tracasseries et le coût associé, même pour des applications de haute sécurité. Les fabricants de scellés, quant à eux, ne sont pas enclins à détailler les contre-mesures qui existent pour protéger leurs produits des attaques : en faisant cela, ils attireraient l'attention sur les défauts de ces produits.

Heureusement, de meilleurs scellés sont possibles. Mais pour faire de meilleurs scellés, il faut comprendre pourquoi les dispositifs existants sont si faciles à neutraliser. Leur talon d'Achille n'est pas la détection de l'accès non autorisé, mais le stockage sûr de la condition d'alarme – c'est-à-dire du fait que l'intrusion a été détectée – jusqu'au moment où le scellé peut être inspecté. Avec les scellés actuels (même les modèles électroniques), il est trop facile pour un escroc de cacher ou d'effacer la condition



6. Un type de scellé électronique nommé *Time Trap* (piège temporel) utilise le principe de l'antipreuve : de l'information, portée par le scellé lui-même, qui indique qu'aucune intrusion n'a eu lieu. Si le scellé détecte une intrusion, il efface cette information. Au moment de la mise en service, le piège temporel sélectionne un nombre aléatoire qui lui servira de clef, et qui n'est connu que de la personne qui met en place le dispositif. Si le scellé détecte une effraction, il utilise l'heure de cet événement et un algorithme nommé fonction de hachage pour calculer un code de sortie spécifique de cet instant. Dans l'exemple illustré ici, le microprocesseur du scellé divise la clef par l'heure, élimine les décimales, puis soustrait 26 à plusieurs reprises du résultat, jusqu'à produire un nombre compris entre 1 et 26, qui correspond à une des lettres de l'alphabet. Cette lettre et l'heure sont alors affichées. Le tout ne prend que quelques microsecondes, après quoi le scellé efface la clef. Si le dispositif n'affiche pas la bonne lettre code au moment de l'ouverture autorisée, c'est que le scellé a subi une effraction.

d'alarme, ou de remplacer le scellé par une contrefaçon neuve qui ne présente aucun signe d'effraction.

Une façon de remédier à cette faiblesse est de prendre le problème à l'envers. Il s'agit, lorsqu'on appose le scellé, de lui incorporer l'information qu'un accès non autorisé n'a pas encore été détecté. Cette information constitue ce que nous appelons une antipreuve. Un tel dispositif peut être mécanique, mais il est le plus souvent électronique. Quand le scellé détecte une intrusion, il efface instantanément son antipreuve. Au moment de l'inspection, l'absence de l'antipreuve indique qu'une intrusion a eu lieu. Avec cette approche, l'agresseur n'a pas à sa disposition une information à cacher ou à effacer. Contrefaire ou copier les composants matériels du scellé ne lui apporte rien s'il ignore quelle antipreuve le scellé stocke. Si l'antipreuve occupe deux octets de mémoire, la probabilité de la deviner est de une chance sur 65 536 ; pour quatre octets, on tombe à une chance sur 4,3 milliards.

Une solution : l'antipreuve

Nous avons développé divers prototypes de scellés à antipreuve. Ils ne coûtent pas grand-chose, sont réutilisables (même les modèles mécaniques) et aucun outil n'est nécessaire pour les apposer ou les retirer. Dans de nombreux cas, les scellés à antipreuve peuvent être placés à l'intérieur comme à l'extérieur du conteneur à surveiller, et beaucoup ne nécessitent pas de lecteur.

Ces scellés à antipreuve permettent également de vérifier automatiquement que l'inspecteur a effectivement vérifié le scellé et ne se contente pas de prétendre qu'il l'a fait, ce qui est un réel problème pour de nombreux programmes de détection d'intrusions. En effet, si l'inspecteur est obligé de déclarer l'antipreuve à ses supérieurs hiérarchiques, l'acte consistant à déterminer l'exactitude de l'antipreuve équivalait à vérifier l'intégrité du scellé.

L'un de nos scellés électroniques à antipreuve est un « piège temporel ». Il est contrôlé par un microprocesseur programmable embarqué. Nous avons assemblé le prototype à partir de pièces revenant à moins de huit dollars au total, un coût qu'une fabrication en masse diminuerait. Si on le place à l'intérieur d'un colis ou d'un conteneur, on évite que la présence de méthodes de détection des intrusions soit visible de l'extérieur, et les intrus éventuels ne vont pas nécessairement remarquer le scellé après s'être introduits. Et avec le scellé apposé à l'intérieur du conteneur, il est plus difficile pour l'intrus d'essayer de faire passer une perturbation du scellé pour un dommage accidentel qui se serait produit durant le transport.

Comment fonctionne notre « piège temporel » ? Quand on l'appose et le met en marche, il choisit un nombre aléatoire en guise de clef. Si le scellé détermine que le conteneur a été ouvert (peu importe si l'intervention est légitime ou non), il active son écran à cristaux liquides. L'écran affiche alors l'heure à laquelle l'ouverture a eu lieu, ainsi qu'un code à deux lettres qui est spécifique de cet instant. Un code différent est attribué à chaque minute ; il est engendré à partir de la clef et d'un algorithme nommé fonction de hachage. Seules les personnes autorisées connaissent le code de hachage pour les instants ultérieurs, et la puce du scellé efface à la fois la formule pour le calculer et la clef en quelques microsecondes dès

que le scellé a détecté une intrusion. Ainsi, les personnes malintentionnées ignorent ce que l'écran devrait afficher au moment où les personnes autorisées ouvriront le conteneur.

Un autre de nos prototypes utilise un lecteur qui communique avec le scellé à l'aide de signaux radio à courte portée. Le scellé est placé à l'intérieur du camion (ou du conteneur) à surveiller. Il détecte en général la lumière, mais il peut utiliser simultanément jusqu'à 14 types de capteurs différents, qui enregistrent le champ magnétique, le mouvement, l'inclinaison, l'accélération, des signatures chimiques, etc.

Pour ce scellé, l'antipreuve est un « slogan » dit par le lecteur électronique. Avant l'ouverture du camion, le lecteur interroge le scellé, et celui-ci renvoie le numéro du slogan qui devrait être dit. Si c'est le bon slogan que l'on entend au moment où le scellé est inspecté, cela signifie qu'aucune intrusion n'a eu lieu. S'il n'y a pas de réponse, ou si ce n'est pas la bonne phrase, cela signifie que quelqu'un s'est introduit auparavant dans le camion, ce qui entraîne l'effacement du numéro du slogan correct.

Le fait d'avoir un code parlé humanise le processus d'inspection et est avantageux d'un point de vue psychologique. Trop souvent, avec les lecteurs automatisés des scellés de haute technologie, les inspecteurs se laissent distraire, se démobilisent mentalement, ne se sentent plus impliqués personnellement et négligent l'examen visuel des détails de la cargaison. Avec des scellés parlants, au contraire, nous avons remarqué que l'auditeur vérifiait quand même attentivement le conteneur et son environnement.

Des recherches insuffisantes

Les scellés font partie de l'expérience humaine depuis fort longtemps. Aussi est-il étonnant que si peu de recherches aient été consacrées à leur amélioration et à la détection d'intrusions. Cette carence est doublement remarquable lorsqu'on sait la facilité avec laquelle les scellés peuvent être déjoués, alors qu'ils jouent un rôle considérable dans la sécurité intérieure, la protection des consommateurs et la lutte contre la prolifération nucléaire. Une situation regrettable, car une meilleure détection des intrusions est à la fois nécessaire et possible. Cela passe par la conception de scellés innovants, mais il ne faut pas oublier que les pratiques de mise en œuvre sont aussi importantes que les scellés eux-mêmes.

Nous remercions la revue *American Scientist* de nous avoir autorisés à publier cet article.

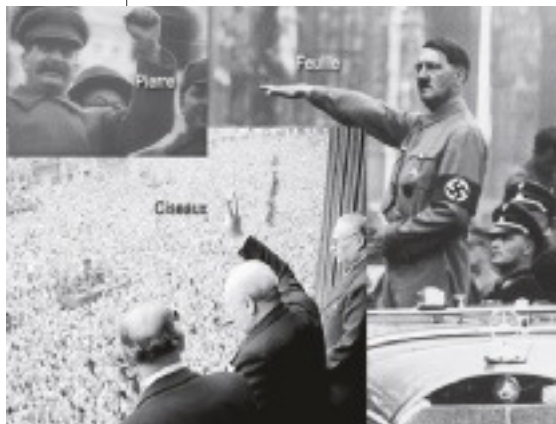
Roger JOHNSTON dirige depuis 1992 l'équipe d'estimation de la vulnérabilité au Laboratoire national de Los Alamos (LANL), aux États-Unis.

Document [traduit en français] du Conseil économique des Nations unies : www.unecp.org/trans/bcf/wp30/documents/d03-13f.pdf

R. G. JOHNSTON, *The « anti-evidence » approach to tamper-detection, in Packaging, Storage & Security of Radioactive Material*, vol. 16, pp. 135-143, 2005.

R. G. JOHNSTON et J. S. WARNER, *The Dr. Who conundrum : Why placing too much faith in technology leads to failure*, in *Security Management*, vol. 49 [9], pp. 112-121, 2005.

D. COLLON (éd.), *7 000 years of seals*, British Museum Press, 1997.



Pierre, feuille, ciseaux...

Joué comme un jeu de conquête, l'ancien jeu chinois de la pierre, de la feuille et des ciseaux éclaire les phénomènes de synchronisation et de maintien de la diversité animale.

Les mathématiciens sont comme les Français : quoi que vous leur disiez, ils le traduisent dans leur propre langue et le transforment en quelque chose de totalement différent.

Goethe

Et de très intéressant... aurait dû ajouter Goethe : ainsi, les mathématiciens ont transformé un jeu d'enfants, *Pierre-feuille-ciseaux*, en un riche modèle des interactions animales. Qui n'a joué à ce jeu ? La pierre (P) casse les ciseaux, les ciseaux (C) découpent la feuille et la feuille (F) enveloppe la pierre. Deux joueurs placent chacun leur main droite dans le dos et choisissent l'un des trois symboles : le poing fermé pour la pierre, la main allongée pour la feuille, deux doigts en forme de ciseaux. Les deux joueurs dévoilent leur choix en avançant la main, parfois en proférant un cri. Si les deux choix sont identiques, ils recommencent, sinon le joueur qui détient le symbole le plus fort des deux a gagné le coup.

Le jeu est ancien et l'on en trouve des traces en Chine dès l'époque Ming (1368-1644). Il est connu dans le monde entier et ses variantes sont innombrables. On y joue avec un plus grand nombre de symboles, on ajoute le puits, le feu, l'eau, etc., ou alors avec d'autres triplets. Dans une version japonaise, les trois symboles sont le guerrier, le tigre et la mère du guerrier. Il y a aussi le feu, le serpent et l'eau, ou encore le serpent, la grenouille et la limace.

Parmi les stratégies du type « jouer P avec la probabilité x , jouer F avec la probabilité y , jouer C avec la probabilité z », nommées *stratégies mixtes*, la stratégie correspondant aux coefficients $x = y = z = 1/3$ est imbattable : si vous l'adoptez, tout joueur sera, à la longue, *ex aequo* avec vous ; il ne pourra ni vous battre... ni même perdre s'il le souhaitait. Cette stratégie ne résoudra pourtant pas toutes les questions que pose le jeu. D'une part, un être humain est incapable d'engendrer dans sa tête des choix vraiment aléatoires avec des probabilités fixées à l'avance : celui qui souhaite jouer la stratégie mixte $1/3-1/3-1/3$ ne le peut pas ! D'autre part, les joueurs humains utilisent l'information sur les coups déjà joués : ils ne peuvent pas s'en empêcher même s'ils essayent.

Le jeu est en partie psychologique et il y a parfois mieux à faire que la stratégie $1/3-1/3-1/3$. Si, par exemple, vous remarquez que votre adversaire n'utilise pas P, alors en jouant invariablement C vous gagnerez toujours (s'il joue F, vous gagnez, s'il joue C, le coup est nul). Bien sûr, il risque de s'apercevoir de votre tactique et de se mettre à jouer P, ce qui vous perdra !

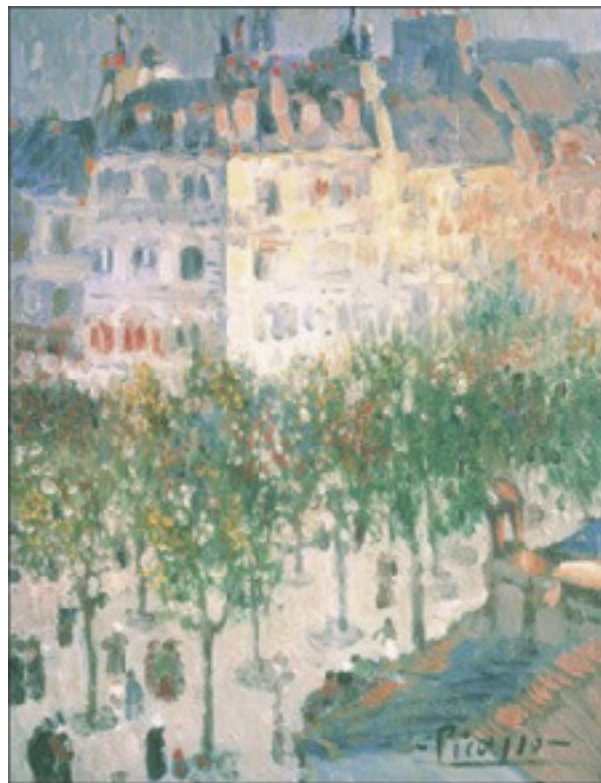
Certains logiciels jouent mieux que les êtres humains, car en repérant les régularités qu'ils adoptent (même quand ils font tout pour ne pas être réguliers), leur attitude donne un léger avantage à une machine attentive et bien programmée. Le programme *Sagace* de Christophe Meyer exploite ainsi les faiblesses humaines.

Comme nous allons le voir dans la suite de cet article, joué sur un graphe, le jeu en apparence simple engendre des comportements collectifs inattendus et délicats à élucider. Certains mécanismes découverts lors de récentes études expliquent les phénomènes de synchronisation et d'uniformisation et éclairent le maintien de la diversité dans le monde biologique.

Les règles du combat

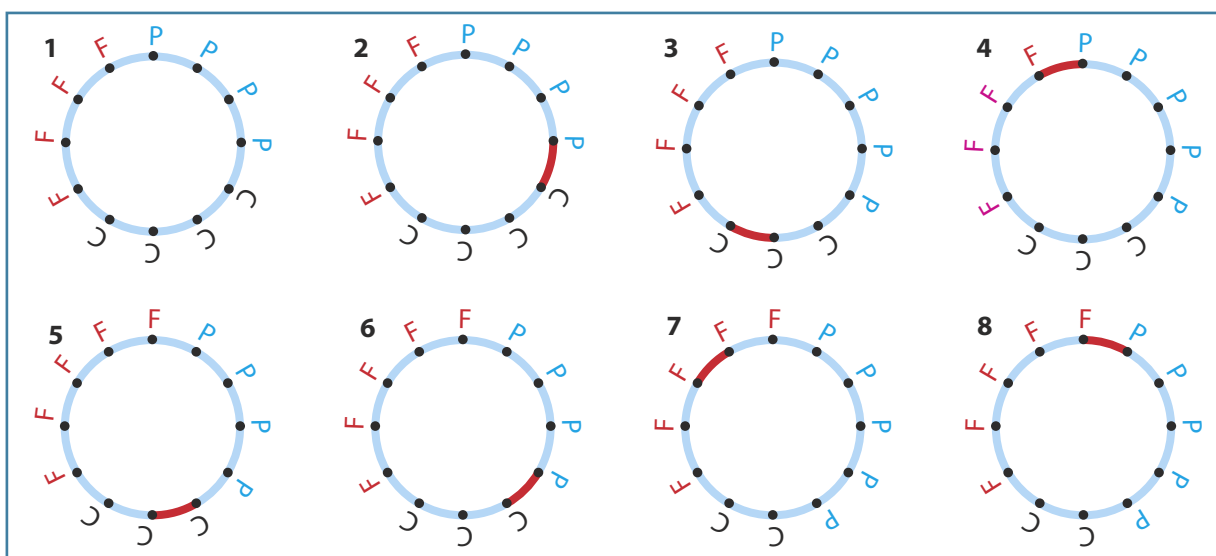
Dans le jeu sur un graphe, un ensemble de points est fixé, les nœuds du graphe. Sur chacun est placé un symbole de type P (pierre), F (feuille) ou C (ciseaux). Certains couples de points sont reliés par les arcs du graphe.

Cette disposition constitue la configuration initiale dont l'évolution se fait selon un processus aléatoire : un nœud N est choisi au hasard équitablement (chaque nœud a la même probabilité d'être choisi) parmi tous les nœuds et un nœud N' directement relié à N est choisi au hasard équitablement aussi. Si les deux symboles portés par N et N' sont de même force, rien ne se produit, sinon le plus faible des deux symboles est remplacé par le plus fort qui occupe alors les deux nœuds N et N'. On recommence, en choisissant un nouveau nœud au hasard, etc. Les combats successifs permettent au gagnant de chaque rencontre d'occuper le terrain de l'adversaire. Ce processus d'évolution est répété un très grand nombre de fois (voir un exemple sur la figure 2).



1. Pierre-feuille-ciseaux est pratiqué sous divers noms : Jan-ken ou Yakyuken au Japon, Kawi-Bawi-Bo en Corée, Jack-En-Poy ou Bato-Bato-Pick aux Philippines, Ko-papir-ollo en Hongrie, etc. Le vainqueur du championnat du monde en 2007 fut l'Américaine Andrea Farina (à gauche). Le jeu, s'il n'est joué qu'une seule fois, est une forme de tirage au sort. Aussi, en 2005, le président de la *Maspro Denkoh Corporation* Takashi Hashiyama, qui ne réussissait pas à choisir si la collection d'œuvres de sa compagnie (comportant des tableaux de Cézanne, Van Gogh et, de Picasso, *Le boulevard de Clichy*, à droite) devait être

vendue par *Christie's* ou *Sotheby's*, organisa une séance de jeu pour fixer sa décision. Les représentants à Tokyo des deux salles des ventes furent conviés à jouer à *Pierre-feuille-ciseaux*. *Christie's* gagna : il avait choisi les ciseaux alors que *Sotheby's* avait opté pour la feuille ! [Pour plus de détails voir le *New York Times* du 29 avril 2005 : <http://www.nytimes.com/2005/04/29/arts/design/29scis.html>]. De toutes les parties de *Pierre-feuille-ciseaux*, c'est sans doute celle à l'enjeu le plus important : la vente portait sur une somme de plus de dix millions de dollars dont plus de 12 pour cent sont revenus à *Christie's*.



2. Le jeu sur un anneau. Un anneau est composé de 12 nœuds chacun lié à ses deux voisins. Au départ, on place quatre P, quatre C et quatre F. Le jeu consiste à tirer au hasard un arc liant deux nœuds, puis à faire jouer l'un contre l'autre les deux symboles (F l'emporte sur P, P l'emporte sur C et C l'emporte sur F). Le symbole qui gagne occupe alors les deux nœuds. Selon le hasard qui fait jouer

ainsi les nœuds les uns contre les autres, soit le jeu préserve la variété, les trois symboles coexistant toujours, soit l'un des symboles finit par disparaître, ce qui entraîne la disparition d'un second symbole, tout l'anneau étant alors occupé par un symbole unique et inamovible (cette occupation unique devient quasi certaine quand le nombre de tirages d'arcs tend vers l'infini).

Dans cet exemple, le combat revient à choisir au hasard deux nœuds voisins de l'anneau, ce qui conduit soit au *statu quo* si les nœuds portent le même symbole, soit à un gain de territoire pour le plus fort. Si tout s'équilibrait, les zones ne feraient que se déplacer, chacune gardant en moyenne quatre éléments et tournant dans le sens des aiguilles d'une montre. Cependant, les tirages au sort des combats ne donnent pas des résultats parfaitement réguliers : les tailles des zones tenues par les C, les P et les F fluctuent et ces fluctuations s'amplifient. Inévitablement, l'une des trois zones disparaît, ce qui entraîne la disparition d'une deuxième zone et il ne reste plus qu'un seul symbole sur tout le terrain. Si la zone qui disparaît en premier est celle des C, les F envahissent tout ; si la zone des P disparaît, les C envahissent l'anneau et si les F disparaissent, les P dominent totalement et définitivement.

Dans le cas d'une configuration initiale irrégulière de C, de F et de P sur un anneau (même beaucoup plus grand) qu'on remplit en tirant au hasard les emplacements des C, des F et des P, l'évolution est de même type. Les déplacements de frontières des différentes zones homogènes amènent assez rapidement une configuration uniforme.

Cette évolution à long terme vers un gagnant définitif (imprévisible) est bien plus générale et se produit pour tout graphe connexe (un graphe connexe est un graphe d'un seul tenant, c'est-à-dire tel que deux nœuds sont toujours reliés par au moins un chemin). Le résultat dont la démonstration est exposée page 94 est : « Pour tout graphe connexe et toute configuration initiale de *pierres*, de *feuilles* et de *ciseaux* sur le graphe, la probabilité pour que le système se retrouve au bout d'un temps fini dans un état uniforme (l'un des symboles occupant tous les nœuds) est de 100 pour cent. »

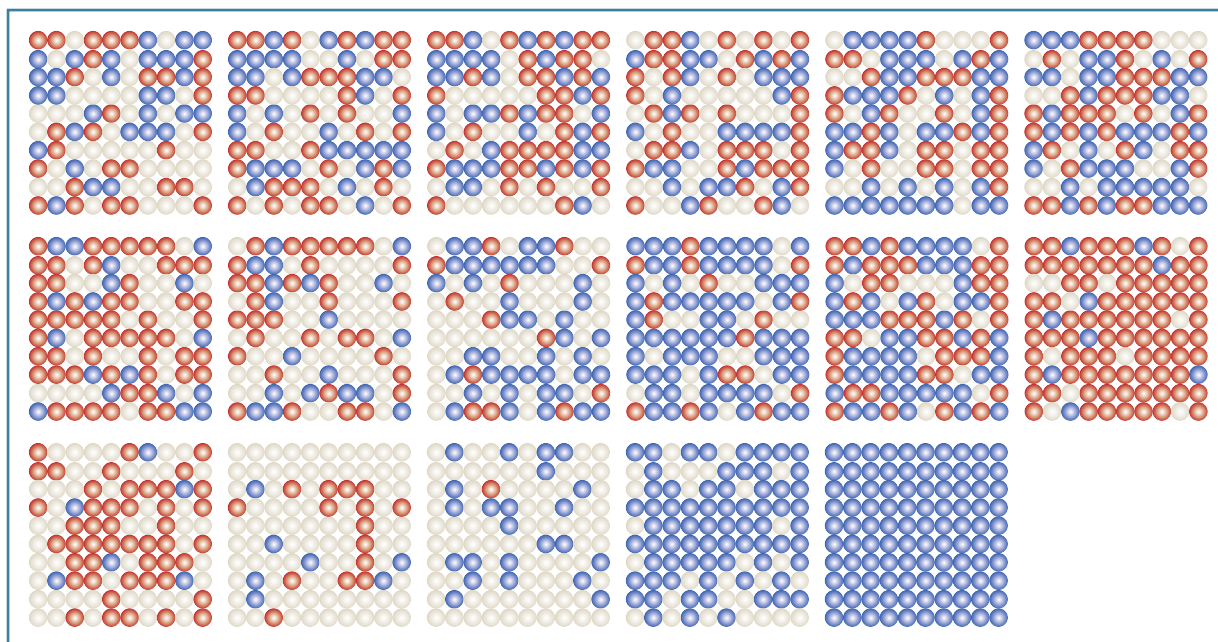
Notons que cela n'interdit pas qu'il existe des suites de tirages particuliers assurant le maintien indéfini des trois symboles sur le graphe (tout comme si vous lancez un dé, vous pouvez obtenir une série ne comportant que des 6). Si l'on considère le graphe circulaire envisagé plus haut et que les arcs sont tirés les uns après les autres en tournant dans le sens des aiguilles d'une montre, il restera toujours des exemplaires des trois symboles. Le sens du résultat d'évolution à long terme est que ce type de suites de tirages préservant la diversité possède une probabilité nulle de se poursuivre à l'infini.

Vérification expérimentale

Nous allons voir que la réalité expérimentale est encore plus intéressante que la théorie mathématique qui indique un peu sèchement qu'au bout d'un certain temps, le système s'uniformise totalement (voir l'encadré 5).

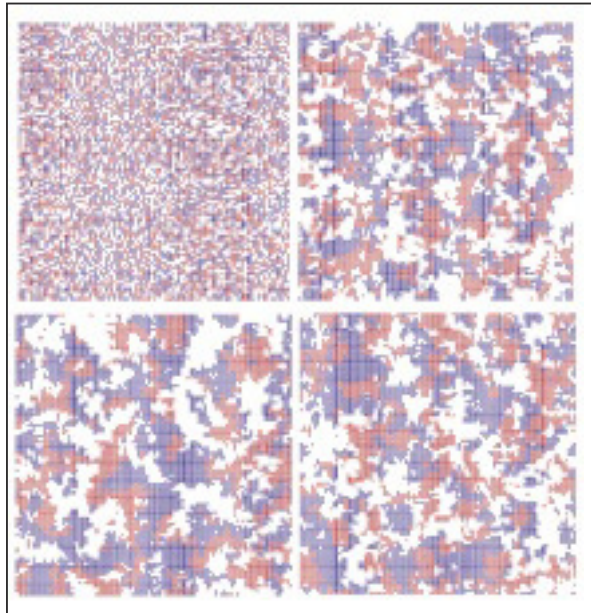
Nous allons d'abord observer l'évolution de graphes complets, c'est-à-dire de graphes dont chaque nœud est relié à chaque autre (graphes ayant donc $n(n-1)/2$ arcs pour n nœuds). De façon à visualiser les changements, nous utilisons une grille carrée où deux cases sont liées qu'elles soient côte à côte sur le dessin ou qu'elles ne le soient pas.

La figure 3 représente l'évolution d'un graphe complet de 100 nœuds. On part d'une configuration aléatoire : rouge pour le symbole *Pierre*, blanc pour *Feuille*, bleu pour *Ciseaux*. Puis on tire successivement 100 arcs du graphe où les combats correspondants sont joués, ce qui donne une nouvelle répartition. Au bout de 17 étapes de 100 combats aléatoires (donc 1 700 combats), le graphe est devenu uniformément bleu : les *Ciseaux* ont gagné. On voit cependant que l'évolution est pas-



3. Le jeu sur un graphe complet (où chaque nœud est relié à chaque autre). Pour visualiser l'évolution, on représente le graphe comme si les nœuds étaient disposés sur un damier. On part d'une configuration aléatoire des trois symboles *Pierre*, *feuille*, *ciseaux* (on utilise pour chacun une couleur, rouge, blanc ou bleu). On représente l'évolution en dessinant le graphe toutes les 100 étapes (100 combats). Dans un premier temps, le graphe reste à peu près homogène avec

33 pour cent de chacun des symboles. Au bout de dix étapes (donc 1 000 combats sur un arc), la couleur bleue domine. Deux cents combats plus tard, le rouge domine. Puis c'est le blanc, puis le bleu, cette fois définitivement. Dans le cas de graphes de taille plus grande, la phase d'uniformisation partielle avec domination cyclique d'un symbole dure plus longtemps avant la victoire définitive d'un des symboles. Les simulations de cette figure ont été réalisées par Rémi Dorat.



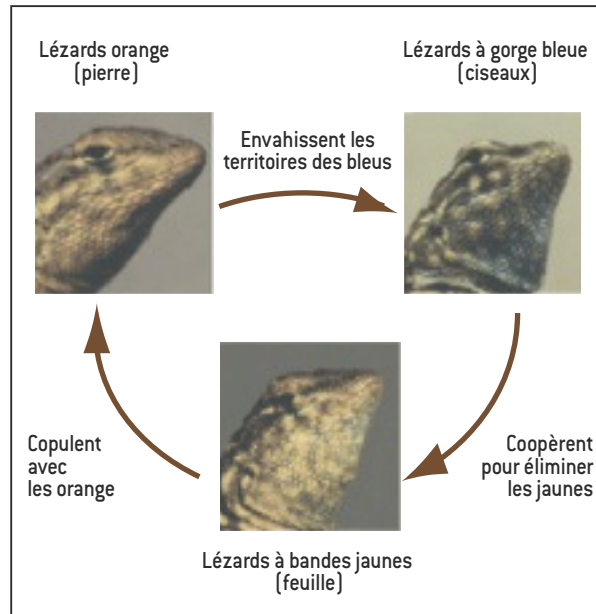
4. Le jeu sur un graphe-damier où chaque case du damier n'est liée qu'à ses huit voisines. Au départ, le damier est rempli aléatoirement par les trois symboles *Pierre*, *feuille*, *ciseaux*. L'évolution fait apparaître des zones homogènes, qui se déplacent tout en gardant une taille à peu près constante. Aucun symbole ne domine. En théorie, la domination totale et définitive d'un symbole se produit au bout d'un

sée par une phase d'oscillation généralisée : sur le dixième dessin, les cases en bleu dominant largement, puis sur le dessin 12, c'est au tour des cases rouges d'être majoritaire, avant de laisser la place au dessin 14 aux cases blanches, avant enfin le retour de la domination bleue, cette fois définitive. Pour des graphes complets plus grands, le phénomène de synchronisation est encore plus net et le temps nécessaire avant son interruption par la domination totale d'un symbole devient long, car la phase de domination tournante possède un très grand nombre d'étapes.

Ce mouvement tournant où chacun des trois symboles prend le dessus à tour de rôle s'explique partiellement : les *Feuilles* vivent en exploitant les *Pierres*, et si les *Feuilles* se mettent à dominer, elles éliminent presque totalement les *Pierres* ; la situation créée est alors favorable aux *Ciseaux* qui n'ont plus à craindre les *Pierres* et qui rencontrent avec une fréquence croissante les *Feuilles*, qu'ils éliminent ; les *Ciseaux* deviennent donc plus nombreux, se multipliant aux dépens des *Feuilles*. Cette domination des *Ciseaux* est favorable aux *Pierres* (plus fort que les *Ciseaux*) qui deviennent alors majoritaires un peu plus tard, etc.

Reste que cette explication est partielle, car si elle justifie que cela tourne dès que l'un des symboles domine, elle n'indique pas pourquoi l'un des symboles prend rapidement le dessus, engendrant l'oscillation synchronisée de tout le graphe : on pourrait imaginer que tout s'équilibre en gros et que les écarts s'effacent vite dès qu'ils se produisent. Nous allons d'ailleurs voir plus loin des graphes d'un autre type où l'équilibre des effectifs est préservé sur des longues périodes.

Dans la nature, un tel phénomène de domination cyclique a été plusieurs fois observé. Le lézard californien, *Uta stansburiana*, étudié par B. Siverno et C. Lively, présente ce



temps fini avec une probabilité de 100 pour cent. En pratique, ce moment vient si tardivement que l'on considère que les trois symboles coexistent dans le monde vivant sont l'une des causes du maintien de la diversité biologique. Le cas du lézard californien *Uta stansburiana* est frappant, car trois variétés de lézards semblent y jouer à *Pierre-feuille-ciseaux*.

comportement de domination cyclique, car il y a trois sortes de mâles qui semblent pris dans un jeu sans fin de *Pierre-feuille-ciseaux* sur un graphe. L'accouplement est à l'origine du cycle car, dans cette cruciale compétition, les mâles à gorge orange l'emportent sur les mâles à gorge bleue qui eux-mêmes dominent les mâles à bandes jaunes qui à leur tour sont plus efficaces sexuellement que les mâles à gorge orange.

Les lézards

Le détail de leur curieux jeu est le suivant. Les mâles à gorge bleue possèdent des harems composés de trois femelles environ et ne cherchent à défendre que de petits territoires. Quand ils sont les plus nombreux, les mâles à gorge orange, qui eux sont particulièrement agressifs, réussissent à s'imposer et font disparaître la quasi-totalité des mâles à gorge bleue. Les lézards orange ont un tempérament dominant à cause d'une forte concentration en testostérone. Cela les conduit d'ailleurs à constituer des harems ayant jusqu'à sept femelles et à exercer leur domination sur de vastes territoires. Ils se substituent donc aux mâles à gorge bleue dont le nombre faiblit rapidement sans toutefois s'annuler. Interviennent alors les mâles à bandes jaunes qui s'introduisent dans les camps des mâles à gorge orange en se faisant passer pour des femelles. Ces mâles à bandes jaunes deviennent de plus en plus nombreux et un moment arrive où ils sont majoritaires. Cela laisse réapparaître les mâles à gorge bleue, car ceux-ci savent reconnaître les mâles travestis et les chassent. Cela conduit au retour de la situation initiale avec une majorité de mâles à gorge bleue. Comme dans le cas de nos graphes complets, un cycle répété



de configurations, où dominent successivement chacune des trois catégories de mâles, se poursuit sans interruption sur les territoires étudiés.

À une échelle bien plus petite, dans le domaine de la chimie, les réactions de Belousov-Zhabotinskii peuvent s'interpréter comme la manifestation d'un jeu de *Pierre-feuille-ciseaux* joué à l'échelle microscopique entre composants chimiques. Dans ce cas cependant, aucun composant ne devient globalement dominant et c'est plutôt à une coexistence des compétiteurs qu'on assiste ; ceux-ci se regroupent par zones homogènes mobiles et créent un jeu animé de dessins étonnants, qu'on a associé à l'idée même de phénomène complexe.

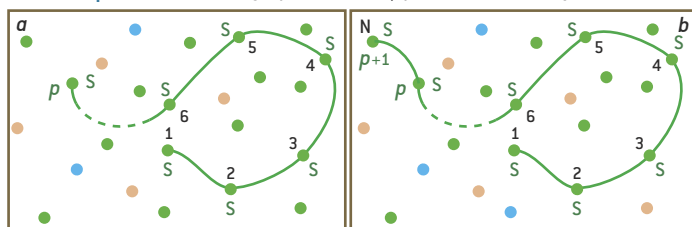
Il est simple d'obtenir l'équivalent de cette dynamique chimique dans le monde abstrait des graphes du jeu *Pierre-feuille-ciseaux* : il suffit de ne pas considérer des graphes complets, mais des graphes réguliers comme ceux correspondant aux

casés d'un damier. Chaque case est reliée à ses huit voisines et non plus à toutes les cases comme précédemment.

La figure 4 montre des configurations typiques observées après un long moment d'évolution lorsque l'on considère des graphes-damiers de ce type. Cette fois, la représentation du graphe sur un damier prend tout son sens, puisque les nœuds du graphe ne sont reliés qu'aux huit nœuds correspondant aux huit cases voisines. Au début, le système est placé dans un état totalement mélangé (1/3-1/3-1/3). Rapidement, apparaissent des zones homogènes des trois couleurs et ces zones se déplacent et changent de forme de manière incessante tout en gardant des caractéristiques globales stables : la taille des zones homogènes est à peu près constante, ainsi que leur vitesse de déplacement et le type de motifs dessinés. Ici, aucun symbole ne domine les autres. L'évolution se traduit par des mou-

5. On peut toujours uniformiser

La démonstration du résultat que l'on aboutit toujours à l'uniformisation se fait en deux étapes. On établit d'abord que : « Pour tout graphe connexe et pour toute configuration initiale donnée, il existe un tirage d'arcs conduisant l'un des symboles à tout envahir. » Le raisonnement se fait par récurrence. Prenons un graphe connexe de taille n , et une configuration initiale donnée. Nous allons montrer que pour tout entier p inférieur ou égal à n , il existe une suite de tirages d'arcs conduisant (quand on effectue dans l'ordre tous les combats déterminés par les arcs) à un paquet de p nœuds rattachés (formant un sous-graphe connexe) portant le même symbole S .



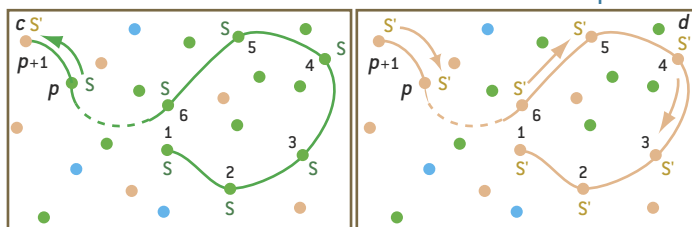
– Cela est évidemment vrai pour $p = 1$.

– Supposons cela vrai pour $p < n$ et considérons la séquence de choix d'arcs qui conduit à un paquet connexe de p symboles S identiques (a). Montrons qu'on peut en déduire une suite de choix d'arcs qui conduit à un paquet connexe de $p + 1$ symboles identiques (d'après le principe du raisonnement par récurrence, cela nous donnera le résultat). Soit un nœud N , voisin de ce paquet : il y en a au moins un, car le graphe entier est connexe. Si ce nœud N porte le symbole S , nous avons un paquet connexe de $p + 1$ symboles S et le raisonnement est terminé (b). Supposons alors que N porte un symbole S' différent de S .

Si S' est battu par S , on tire l'arc liant le nœud N au paquet de S et l'on a une suite de choix d'arcs (c) qui conduit à un paquet connexe de S de taille $p + 1$: le raisonnement est terminé.

Si S est battu par S' , on choisit la même séquence d'arcs que celle qui conduisait au paquet de p symboles S , suivie d'une suite de tirages progressifs (d) de ces arcs de proche en proche qui font que S' envahit tout le paquet de S , lequel devient donc un paquet connexe de $p + 1$ symboles S' . Cela est possible, car le paquet de p symboles S est connexe. On a alors un paquet de $p + 1$ symboles S' . Dans chaque cas possible, on dispose d'une

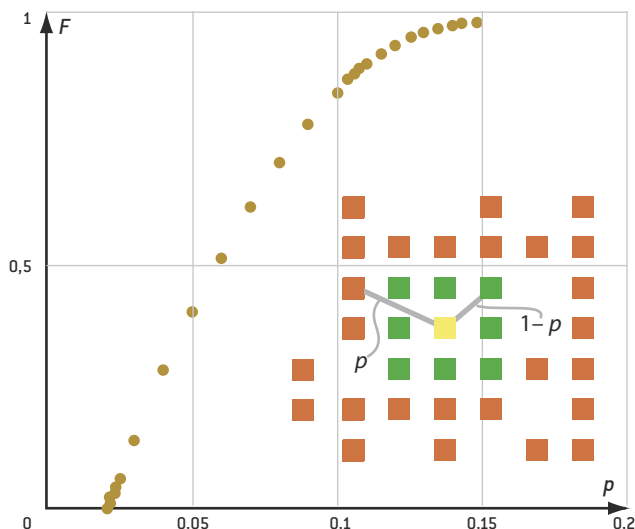
suite des choix d'arcs conduisant à un paquet connexe de $p + 1$ symboles identiques. Le raisonnement par récurrence est terminé.



Notons (c'est important pour la seconde partie du raisonnement) que la longueur de la suite d'arcs qui fait passer d'un paquet de p nœuds identiques à un paquet de $p + 1$ nœuds identiques s'obtient en ajoutant au plus p arcs à celle nécessaire pour avoir un paquet de p nœuds identiques. Au total, si le graphe possède n nœuds, on sera certain d'avoir un graphe uniforme en utilisant une suite d'arcs déterminant les combats successifs de longueur au plus $1 + 2 + 3 + 4 + \dots + (n - 1) = n(n - 1)/2$.

Montrons maintenant que l'uniformisation survient avec une probabilité de 100 pour cent. Cette seconde partie du raisonnement utilise l'évidence probabiliste qu'en répétant un tirage au hasard dans l'attente d'un événement ayant une probabilité non nulle de se produire, il finit par se réaliser, ou plus exactement, il se réalise « à l'infini » avec une probabilité de 100 pour cent (en lançant un dé indéfiniment, vous finirez par tomber sur un 6 : s'il est possible, en théorie, de ne jamais avoir de 6, cet événement a une probabilité nulle : si l'on attend longtemps, on est certain d'avoir un 6).

Partant d'un graphe de n nœuds dans une configuration donnée, nous savons qu'il existe une suite d'arcs qui uniformise le graphe. Cette suite a une longueur inférieure à $n(n - 1)/2$. La probabilité qu'un arc donné soit choisi lors d'une étape du processus d'évolution aléatoire est plus grande que $1/(nk)$ (un nœud est choisi avec une probabilité $1/n$ et un voisin avec une probabilité $1/k$ au plus, où k est le nombre maximum de voisins que possède un nœud du graphe). La probabilité pour qu'une suite d'arcs qui uniformise le graphe soit choisie est donc supérieure à $(1/(nk))^{n(n-1)/2}$, un nombre strictement positif. À l'infini, la probabilité pour qu'au moins une fois une séquence conduisant à l'uniformisation soit choisie est donc de 100 pour cent.



vements ininterrompus des frontières interzones. La zone de séparation entre *Pierres* et *Feuilles* se déplace, car les *Pierres* reculent devant les *Feuilles* ; de même, les *Feuilles* reculent devant les *Ciseaux* et les *Ciseaux* devant les *Pierres*. L'examen détaillé des effectifs montre une oscillation limitée autour de 1/3 de la proportion de chacun des symboles. Cette fluctuation apparaît comme une variation statistique inévitable qui, ne prenant pas d'ampleur, ne conduira en pratique jamais à la domination d'un symbole.

Temps de convergence différents

Les deux cas rencontrés, oscillation synchronisée de tout le graphe ou formation de zones homogènes mobiles, correspondent à deux modes de la dynamique d'évolution des graphes de *Pierre-feuille-ciseaux*. Dans les deux types de graphes, nous savons qu'à long terme l'un des états envahit tout. Cependant, le temps nécessaire pour que se produise cette domination définitive est très différent dans les deux cas. Dans le cas des graphes complets, ce temps augmente lentement en fonction de la taille du graphe. Dans le cas des graphes-damiers, l'instant de la stabilisation est repoussé bien plus loin, si bien qu'en pratique, le système reste dans un état mixte où les trois symboles coexistent.

La théorie nous indique pourtant que l'un des états finit toujours par dominer ; la simulation nous montre qu'avec les graphes complets, cette domination arrive rapidement, alors qu'avec les graphes-damiers, sa venue est tellement éloignée dans le temps que c'est une stabilisation avec coexistence des trois états qu'il faut considérer comme l'évolution normale à moyen terme de la dynamique du graphe. Une étude détaillée des caractéristiques de ces dynamiques a été menée par Rémi Dorat à coup de simulations massives et prolongées. Il montre que le temps d'uniformisation croît exponentiellement dans le cas des graphes-damiers.

La situation de coexistence des trois symboles est sans doute une explication du maintien de la diversité dans certaines situations biologiques. Le terrain où trois espèces se rencontrent (les dominations se conformant à un schéma circulaire du type *Pierre-feuille-ciseaux*) est assimilable à un graphe de type graphe-damier. La dynamique de domination montre alors

6. Jeux intermédiaires *Pierre-feuille-ciseaux*. On définit une famille continue entre le jeu sur un graphe complet et le jeu sur un graphe-damier. On se fixe un paramètre p entre 0 et 1. À chaque coup, on choisit un premier nœud N au hasard uniformément parmi tous les nœuds. Pour le second nœud N' , on choisit entre deux méthodes A et B, la méthode A choisie avec une probabilité p et la méthode B avec une probabilité $1 - p$. Si la méthode A est choisie on prend pour N' un nœud quelconque du graphe. Si la méthode B est retenue, on choisit N' parmi les huit voisins de N . Les nœuds N et N' ainsi déterminés livrent alors combat. Pour $p = 0$, tout se passe comme quand on jouait au jeu sur un graphe-damier, pour $p = 1$, comme si on jouait sur un graphe complet. Pour p assez petit (jusqu'à 0,002), l'évolution du graphe se fait en préservant un bon mélange : les trois symboles restent présents en proportion équilibrée 1/3-1/3-1/3. Lorsque p se trouve entre 0,002 et 0,15, la proportion de chaque symbole fluctue dans le système, ce que l'on mesure par un paramètre de fluctuation F ($F = 0$, pas de fluctuation ; $F = 1$, variation maximale des effectifs). Au-delà de 0,15, le système devient violemment oscillant comme dans le cas d'un graphe complet : chacun des symboles domine à tour de rôle, le graphe passant par des états où un symbole occupe la quasi-totalité du graphe, toutes les cases du graphe sont synchronisées (ces simulations ont été menées par A. Szolnoki et G. Szabó).

sur le terrain des motifs plus ou moins semblables à ceux vus lors des simulations.

Benjamin Kerr, Margaret Riley, Marcus Feldman et J. M. Bohannan, dans un article de la revue *Nature*, proposent une telle interprétation de leur expérience réalisée avec trois souches d'*Escherichia coli*, en relation non transitive comme les symboles *Pierre-feuille-ciseaux*. La conclusion de leur étude est conforme à ce qui est montré par les simulations. Lorsque l'espace où est réalisée l'expérience est petit et, plus généralement, lorsque les relations entre souches sont assimilables à celles décrites par un graphe complet, l'une des souches domine rapidement. En revanche, lorsque les interactions se déroulent sur un espace plus grand, cette fois assimilable à un graphe de type damier, les trois souches coexistent avec des mouvements des zones occupées par chacune d'elles, mouvements assez semblables à ceux vus lors des simulations.

Des études sur des graphes intermédiaires entre les graphes complets et les graphes damiers ont été réalisées par György Szabó, Attila Szolnoki et Rudolf Izsák. Divers phénomènes de transition de phases ont pu être mis en évidence, et ils permettent de mieux comprendre toutes ces dynamiques discrètes nées d'un jeu élémentaire (voir la figure 6). La science n'a jamais fini d'étudier les jeux : même les plus simples suggèrent des modèles explicatifs utiles.

Jean-Paul DELAHAYE est professeur d'informatique à l'Université de Lille.

R. DORAT, *Étude de la domination cyclique sur les graphes*. Publication du Laboratoire d'Informatique Fondamentale de Lille, 2007.

C. MEYER, *Sagace* : <http://ch.meyer.free.fr/favorite.htm>

G. SZABÓ, A. SZOLNOKI, et R. IZSÁK, *Rock-scissors-paper game on regular small-world networks*, in *J. Phys. A: Math. Gen.*, 37, pp.2599-2609, 2004.

B. SIVERNO, C.M. LIVELY, *The rock-paper-scissors game and the evolution of alternative male reproduction strategies*, in *Nature*, vol. 380, pp. 240-243, 1996.

B. KERR, M. RILEY, M. FELDMAN et J. M. BOHANNAN, *Local dispersal promotes biodiversity in a real-life game of rock-paper-scissors*, in *Nature*, vol. 418, pp. 171-174, 11 juillet 2002.

A. SZOLNOKI, G. SZABÓ, *Phase transitions for rock-scissors-paper game on different networks*, in *Phys. Rev. E* 70, 2004 : <http://www.mfa.kfki.hu/~szabo/szabopub.html>

La petite sirène d'Ambroise Paré

Le chirurgien français a décrit un « monstre »,
un être humain doté, entre autres étrangetés d'ailes,
d'une corne et d'un seul pied... d'oiseau. Jusqu'où était-il réaliste ?

Dans son traité *Des Monstres et Prodiges* (1575), le chirurgien français Ambroise Paré décrit un « Monstre merveilleux » : « Du temps que le pape Jules II suscita tant de malheurs en Italie et qu'il eut la guerre contre le roi Louis XII (1512) laquelle fut suivie d'une sanglante bataille donnée près de Ravenne : peu de temps après on vit naître en la même ville un monstre ayant une corne à la tête, deux ailes et un seul pied semblable à celui d'un oiseau de proie ; à la jointure du genou [était] un œil ; et participant de la nature de mâle et de femelle, comme tu vois par ce portrait (voir page ci-contre). » C'est une sirène, car on oublie souvent que ces créatures sont avant tout des humains affublés d'un membre inférieur d'oiseau (la tradition d'une queue de poisson est plus tardive)...

À l'époque d'Homère, il s'agit d'êtres monstrueux : des femmes dans la partie supérieure du corps et oiseaux des hanches aux pieds. On peut rapprocher ces figures mythiques des observations ethnologiques et anthropologiques de Pliny l'Ancien, qui nous parle du peuple des Monocoli vivant près du pays des Troglodytes (l'Inde ou l'Éthiopie, selon les auteurs). Ils étaient, à l'en croire son *Histoire Naturelle*, « affublés d'un seul membre inférieur médian avec lequel ils se déplaçaient avec agilité ». Ce peuple était également nommé les Sciapodes (en grec « pieds d'ombre »), car ces individus auraient été dotés d'un pied énorme dont l'ombre les protégeait, aux plus chaudes heures de l'été, lorsqu'ils s'allongeaient sur le dos, la jambe relevée.

Dans l'iconographie gréco-romaine, une patte d'oiseau remplace la moitié inférieure du corps. En cela, cette iconographie est conforme à l'observation de sirène faite par Ambroise Paré. Dans tous les cas, aucune tradition antique ne parle de sirène sous la forme de femme-poisson ; il s'agirait d'une perversion de la légende apparue dès le Moyen Âge. Comment expliquer les individus décrits par les auteurs antiques et dénommés en référence à ces créatures mythologiques ?

Les embryologistes parlent d'une dysgénésie caudale, c'est-à-dire un défaut de développement de la partie inférieure de l'embryon à un stade très précoce de son développement, dès la troisième semaine de gestation : ce peut être soit un défaut de migration du mésoderme (l'un des trois feuilletts embryonnaires) vers l'extrémité inférieure, soit une nécrose partielle et très localisée d'un tissu embryon-

naire, le mésoblaste postérieur (on parle alors de syndrome de régression caudale). Selon la gravité de cette anomalie, les membres inférieurs, l'appareil génito-urinaire et les vertèbres lombo-sacrées sont atteints. Dans le pire des cas, ils sont tous les trois réduits à l'état vestigial : l'anus n'existe pas, les vertèbres sacrées et le coccyx ne sont pas aboutis, les membres inférieurs sont grêles et fusionnés, les reins sont mal formés et les organes génitaux externes sont absents. Le bassin est souvent étroit, fermé par une fusion pubienne à l'emplacement du périnée.

Dans les formes ultimes de cette maladie, les cotyles (les cavités semi-circulaires de l'os de la hanche où se logent les fémurs) se réunissent en arrière dans le prolongement du rachis mal formé, entraînant une fusion des ébauches des membres inférieurs par leurs bords externes. La fréquence de cette malformation congénitale gravissime est d'un cas pour 60 000 naissances.

On distingue trois degrés de malformation : la sirénomélie, où les nouveau-nés sont dépourvus de pied ; l'uro-mélie (ils n'ont qu'un pied) ; la symélie, avec un membre unique doté de deux pieds, le plus souvent fusionnés (huit-dix orteils), en flexion permanente à 90 degrés sur le tronc, les rotules placées latéralement. Toutes ces malformations, et surtout l'atteinte rénale, sont incompatibles avec la vie, et l'enfant est bien souvent mort-né.

Dans l'iconographie de la sirène d'Ambroise Paré, les deux membres inférieurs sont fusionnés en un membre unique terminé par un pied griffu d'oiseau. En réalité, les syndactylies (fusions congénitales d'orteils ou de doigts) sont fréquentes chez les individus syméliques, d'où des orteils de taille et de forme anormales, en nombre excessif pour un seul membre. De telles anomalies anatomiques expliqueraient cet aspect propre aux pieds des sirènes, la description anatomique qui en est faite ou l'analogie avec les serres des oiseaux. En revanche, les ailes, la corne et l'œil ne sont qu'une exagération « artistique » ajoutée pour grossir un tableau tératologique déjà bien fourni.

Philippe CHARLIER est médecin-légiste et paléopathologiste à l'Hôpital Raymond Poincaré, à Garches (AP-HP, Univ. Versailles-Saint-Quentin).

Ph. CHARLIER (dir.), *1^{er} Colloque international de pathographie* [Loches, avril 2005], Paris, De Boccard, 2006.



Pont de Tacoma : la contre-enquête

Nos limiers physiciens déclarent que l'effondrement du pont de Tacoma, en 1940 aux États-Unis, n'est pas dû à un phénomène de résonance.

Le 7 novembre 1940, après d'impressionnantes oscillations, le pont de Tacoma s'écroule, sous les yeux médusés de nombreux témoins. Cette catastrophe sans victimes, mais abondamment filmée et photographiée, a acquis depuis une renommée mondiale, avec un coupable tout désigné : le phénomène de résonance. L'erreur judiciaire pointe car, contrairement à ce qui est affirmé dans nombre d'ouvrages, l'accusé n'est pas à l'origine de la catastrophe.

Commençons l'enquête par les faits. Avant les funestes événements du 7 novembre 1940, la victime était déjà bien connue des services de sécurité de l'État de Washington. Dès sa prime jeunesse, le pont de Tacoma, inauguré le 1^{er} juillet 1940, se comportait dangereusement. Son tablier avait la fâcheuse tendance d'osciller verticalement dès qu'une petite brise se levait. Ces oscillations d'une amplitude de plusieurs dizaines de centimètres étaient impressionnantes. Mais l'élasticité naturelle des matériaux permettait au pont de se déformer sans rompre. Soumis à des sollicitations extérieures telles que le vent ou le passage des véhicules, le pont se mettait à osciller selon des modes bien définis. Ses concepteurs estimaient toutefois que ces oscillations resteraient limitées et qu'il n'y avait pas de soucis majeurs de sécurité. Le début de l'automne sembla leur donner rai-

son : le pont résista aisément à des vents de plus de 80 kilomètres par heure.

Mais le matin du 7 novembre, sous l'effet de vents modérés (60 à 70 kilomètres par heure), le tablier du pont se met à vibrer plus que d'habitude, avec une amplitude de plus de un mètre, assez pour interdire le trafic. À dix heures du matin, coup de théâtre : les oscillations verticales se transforment en torsions périodiques ! À partir de ce moment, à chaque va-et-vient, l'inclinaison du tablier par rapport à l'horizontale augmente, jusqu'à atteindre des déplacements du tablier de neuf mètres et une inclinaison de 45 degrés. Ce nouveau régime d'oscillation est fatal : au bout d'une demi-heure, des morceaux du pont commencent à tomber dans le fleuve, et à 11 h 02, près de 200 mètres de tablier se détachent. Le pont de Tacoma a vécu.

La résonance dispose d'un alibi

À quoi attribuer ce désastre ? La rumeur ne tarda pas à accuser le phénomène de résonance et, encore de nos jours, c'est cette fausse explication qui est colportée. Rappelons ce qu'est une résonance avec l'exemple d'un enfant sur une balançoire. À cause des frottements, les oscillations de la balançoire cessent d'elles-mêmes si l'on n'apporte pas continuellement de l'énergie au système. Les pertes d'énergie étant faibles, il



1. Résonance avec une balançoire : une faible poussée, imprimée à la même fréquence que les oscillations de la balançoire, compense les pertes d'énergie par frottements et augmente peu à peu l'amplitude du mouvement.



2. Avant de s'effondrer, le tablier du pont de Tacoma a subi, sous l'influence de vents modérés, des oscillations verticales ainsi que des mouvements périodiques de torsion le long de son axe, l'inclinaison allant jusqu'à 45 degrés.

suffit à un parent compatissant de pousser régulièrement l'enfant pour entretenir le mouvement ou pour en augmenter l'amplitude. S'il règle précisément la fréquence de ses interventions sur le rythme des oscillations de la balançoire, il lui suffira d'exercer une toute petite force pour obtenir un effet très important : c'est le phénomène de résonance.

Une structure complexe (pont, immeuble, véhicule, etc.) présente en général différents modes de vibration, chacun pouvant entrer en résonance à une fréquence caractéristique. L'effondrement partiel d'une autoroute californienne, en 1989, à la suite d'un séisme, semble ainsi dû à une résonance à deux hertz, l'une des fréquences propres de vibration de la structure autoroutière : le terrain sous-jacent filtrait les fréquences des mouvements du sol et transmettait justement celles qui étaient proches de deux hertz.

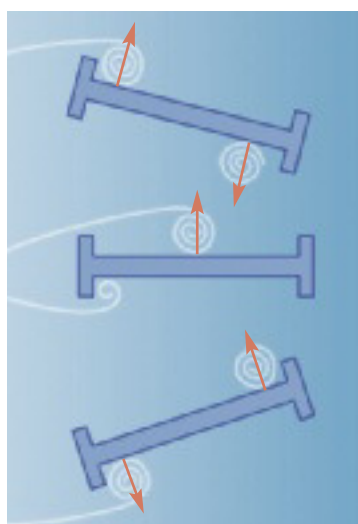
En était-il de même pour le pont de Tacoma, le vent turbulent jouant le rôle du tremblement de terre ? La réponse est « oui » pour les oscillations que connaissait le pont depuis sa construction, mais « non » pour celles qui ont provoqué sa rupture. Les rafales de vent ne pouvaient avoir l'énergie et la régularité nécessaires sur une durée aussi longue que celle observée : plus d'une demi-heure, avec une période de cinq secondes, c'est-à-dire une fréquence de 0,2 hertz.

L'accusation ne se laissa pas démonter. Faisant preuve d'une grande culture physique, elle fit appel aux « allées de von Karman ». Il s'agit de rangées périodiques de tourbillons semblables à celles qu'on observe dans l'eau derrière les piliers des ponts. Elles apparaissent quand un fluide s'écoule assez vite autour d'un obstacle fixe. À la série de tourbillons aérodynamiques créés par le tablier du pont est ainsi associée une dépression qui s'exerce alternativement sur le dessus puis le dessous du tablier.

La force périodique qui s'ensuit aurait-elle fait entrer le pont de Tacoma en résonance ? Le réquisitoire est séduisant, mais il tourne court. En effet, le calcul de la fréquence d'émission de ces tourbillons donne un hertz, soit cinq fois plus que le 0,2 hertz observé par les témoins ou mesuré en soufflerie sur des maquettes.

Quel est alors le vrai coupable ? On ne peut décrypter le comportement du pont de Tacoma en faisant intervenir séparément les oscillations du pont et l'écoulement aérodynamique. En revanche, tout s'éclaire quand on comprend que ces deux phénomènes se couplent et s'amplifient l'un l'autre.

La section du tablier a une forme de H très aplati : les barres verticales correspondent aux parapets qui se trouvent de part et d'autre de la chaussée. Lorsque le vent souffle de côté, il rencontre un parapet. Cet obstacle vertical crée naturellement deux tourbillons identiques de part et d'autre du tablier. Leurs effets mécaniques sur le pont se compensent. Cette symétrie est cependant brisée dès que la chaussée s'incline : le tourbillon qui se trouve sur la face du pont protégée du vent grossit et tire le tablier vers le haut, ce qui accentue l'inclinaison. En revanche, quand le mouvement s'inverse, le tourbillon décroche du parapet et parcourt la largeur du tablier avant de s'éloigner tandis que, de l'autre côté, un autre tourbillon croît.



3. Quand le vent souffle de côté (ici de gauche à droite), les parapets du pont engendrent des tourbillons d'air au-dessus et au-dessous du tablier. Si le tablier reste horizontal, les forces (flèches rouges) exercées sur le pont par les tourbillons se compensent. En revanche, si le tablier oscille périodiquement autour de son axe, les tourbillons amplifient ces mouvements de torsion.

Un complot torsion-tourbillons

L'oscillation du tablier crée donc des tourbillons synchrones, dont l'action amplifie l'oscillation ; en retour, celle-ci donne naissance à des tourbillons encore plus importants, et ainsi de suite. C'est ce mécanisme subtil qui a provoqué *in fine* l'effondrement du pont de Tacoma.

Ainsi, il ne s'agit pas d'une résonance sous l'effet d'une force périodique. Telle une anche oscillant sous le souffle régulier du clarinettiste, le pont de Tacoma est entré en vibration sous l'effet du vent. Ces vibrations se sont couplées à des tourbillons, couplage qui a amplifié les mouvements.

Les procureurs très pointilleux pourraient souligner une faiblesse de notre argumentation. La reconstitution laisse en effet de côté la naissance des mouvements de torsion. Ce point est toujours controversé : rupture ou déplacement subit d'un câble de soutien ? Couplage entre différents types de tourbillons ? Comme dans toutes les bonnes affaires criminelles, il reste une petite part de mystère.

Jean-Michel COURTY et Édouard KIERLIK sont professeurs de physique à l'Université Pierre et Marie Curie, à Paris.

D. GREEN et W. G. UNRUH, *The failure of the Tacoma bridge : a physical model*, in *American Journal of Physics*, vol. 74, n° 8, pp. 706-716, 2006.

Y. K. BILLAH et R. H. SCANLAN, *Resonance, Tacoma Narrows bridge failure, and undergraduate physics textbooks*, in *American Journal of Physics*, vol. 59, n° 2, pp. 118-124, 1991.

<http://www.youtube.com/watch?v=AsCBK-fRNRk>

Zoologie

Bernard Heuvelmans

Un rebelle de la science

Jean-Jacques Barloy (199 pages, 25 euros),
L'Œil du Sphinx, 2007.

Les félins encore inconnus d'Afrique

Bernard Heuvelmans (290 pages, 28 euros),
L'Œil du Sphinx, 2007.

Bernard Heuvelmans (1916-2001) se fit connaître du grand public en 1955, avec la publication de son livre *Sur la piste des bêtes ignorées*, qui fit de lui le fondateur de la cryptozoologie, la science des « animaux cachés », à savoir ceux qui n'ont pas encore été répertoriés par la science. Les cryptozoologues du monde entier voient en lui un père spirituel, même si certains ont une fâcheuse tendance à s'écarter des principes qu'il s'efforça de respecter dans ses investigations. De ce fait, il n'est guère possible de vraiment comprendre la cryptozoologie, son attrait, ses réussites et ses faiblesses, sans savoir qui fut Heuvelmans. La biographie très intimiste et pleine d'admiration que lui consacre Jean-Jacques Barloy révèle un personnage étonnant, trompettiste de jazz dans la Belgique d'avant-guerre, soutenant une thèse de zoologie à Bruxelles en 1939, menant la vie de bohème à Saint-Germain-des-Prés durant l'après-guerre tout en publiant sur les thèmes les plus variés, pratiquant le naturisme à l'île du Levant (Var) dans les années 1950... Mais c'est la quête des bêtes ignorées, du Yéti au Grand-Serpent-de-Mer, qui lui conféra la notoriété. Fut-il un rebelle de la science, comme le proclame le titre du livre de J.-J. Barloy ? Il n'occupa jamais de poste académique et certaines de ses idées et interprétations, dans le domaine si particulier de la cryptozoologie, purent choquer certains esprits timides. Mais, comme le rappelle dans sa préface le peintre Arika Lindbergh, qui le connut mieux que personne, Heuvelmans n'a pas été rejeté par la « science officielle ». Son anticonformisme ne l'empêcha nullement de se plier aux principes de la recherche scientifique, et en cela il s'oppose totalement à certains illuminés qui mélangent allègrement bêtes ignorées, extraterrestres et apparitions surnaturelles. Dans sa

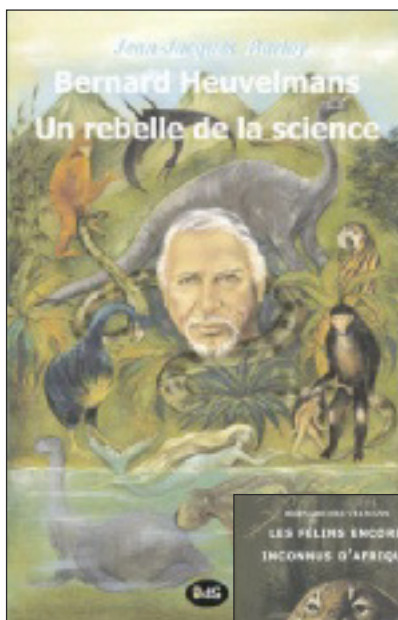
postface, le sociologue des sciences Benoît Grison analyse l'approche de Heuvelmans et rappelle divers succès récents de la cryptozoologie, du Vietnam à l'Amazonie en passant par l'océan Indien, où l'on découvre encore des espèces de grands vertébrés jusque-là inconnues de la science.

Pour bien comprendre comment Heuvelmans concevait l'enquête cryptozoologique, on lira avec intérêt et plaisir son ouvrage posthume sur les félins inconnus d'Afrique. Un don remarquable pour dénicher les témoignages les plus obscurs, une grande culture zoologique et une audace certaine dans la formulation des hypothèses, voilà ce qui caractérise notre cryptozoologue. Dans son style élégant, il promène le lecteur aux quatre coins de l'Afrique, avec une gradation savante dans la nature des problèmes abordés. Que des sous-espèces, voire des espèces, non encore répertoriées de lions, de panthères ou de guépards, puissent encore se dissimuler dans des régions peu accessibles du

continent, le lecteur pourra sans doute l'accepter aisément, sachant toutefois, comme le rappelle justement l'auteur, que la variation individuelle ou géographique peut être trompeuse. Quand Heuvelmans, pour expliquer certains témoignages bien mystérieux, postule la survie en Afrique de félins à dents en sabre, que l'on ne connaît actuellement qu'à l'état fossile, et de plus leur prête des mœurs aquatiques ou montagnardes plutôt inattendues, on le suit avec plus d'hésitation, même si les interprétations sont ingénieuses. Comme souvent en cryptozoologie, on peut se demander à quel point les témoignages sont fiables et si des influences culturelles n'ont pas joué sur la perception des témoins. On notera aussi que beaucoup de rapports concernant ces félins ignorés remontent à la période coloniale. Dans l'Afrique d'aujourd'hui, plus moderne, émancipée mais déchirée de conflits, suivre la piste des bêtes ignorées est peut-être encore plus difficile qu'au temps des administrateurs coloniaux et des chasseurs de fauves. Quoi qu'il en

soit, l'œuvre de Heuvelmans stimule l'imagination scientifique et fait aussi rêver, avec un peu de nostalgie pour l'époque où compétitions sportives et raids écologiques n'avaient pas encore supplanté les voyages d'exploration. Il faut donc féliciter les éditions de L'Œil du Sphinx d'avoir ainsi entrepris la publication des nombreux inédits laissés par Heuvelmans ; les suivants sont attendus avec impatience.

*Eric Buffetaut, CNRS,
Laboratoire de géologie de l'ENS, Paris*



Préhistoire

L'Atlas des origines de l'homme

Une histoire illustrée

Douglas Palmer (192 pages, 36 euros),
Delachaux et Niestlé, 2007.

La Préhistoire ne se résume ni à la France ni au continent européen. Les nombreuses phases de notre évolution ne se sont pas déroulées de manière continue au même endroit, mais souvent en plusieurs

régions différentes et de façon concomitante. Voilà une réalité que les vulgarisateurs de la Préhistoire ont le plus grand mal à faire entendre au lecteur français ! Pour résoudre cette difficulté, la présentation en atlas est un excellent moyen d'entraîner le lecteur sur plusieurs terrains à la fois, par une suite de fiches synthétiques. De ce point de vue, l'ouvrage est particulièrement pédagogique. De plus, il exprime un point de vue anglo-saxon, qui peut parfois nous choquer, mais qui est extrêmement rafraîchissant. D'abord parce que l'auteur ne présente pas les thèmes classiques dans le même ordre ni de la même façon que les ouvrages de synthèse auxquels nous sommes habitués. Il n'a pas non plus les mêmes précautions de langage.

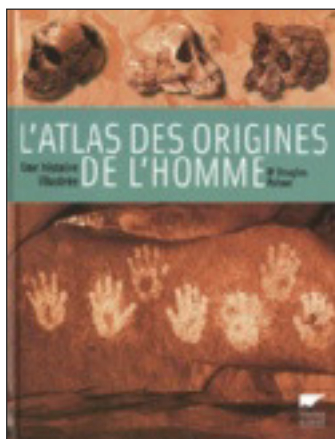
Le livre se décompose en trois parties. La première plante le décor. D. Palmer y résume à grands traits toutes les nouvelles découvertes. Il brosse un portrait rapide de l'évolu-

tion humaine tout en réfutant, sans en avoir l'air, certains préjugés racistes, comme celui de la couleur de la peau.

La deuxième partie est plus classique : D. Palmer présente les grandes familles de l'évolution humaine, mais, contrairement à l'usage, en commençant par les fossiles les plus récents. Ce parti pris pourrait se révéler dangereux, en sous-entendant que l'auteur cherche à démontrer une évolution orientée vers une fin glorieuse : nous. Heureusement, il n'en est rien et D. Palmer prépare simplement le lecteur à la discussion finale : qu'est-ce qu'un homme ? À partir de quand peut-on dire qu'un fossile appartient au genre *Homo* ? À chaque fois, l'auteur présente un bref historique des découvertes des fossiles et de tous les problèmes qui se posent à leur rencontre (la chronologie de l'occupation de l'Europe, par exemple).

La troisième partie est militante. Il s'agit de présenter la dispersion de notre espèce hors d'Afrique, depuis 100 000 ans. Et de démontrer, une fois encore, sa profonde unité, malgré certaines différences présentées comme superficielles. À travers une multitude d'exemples bien documentés et illustrés, toujours situés sur une carte, D. Palmer recadre le lecteur et pose les bonnes questions. Des encadrés colorés proposent une approche plus pointue. Des vignettes résument le contenu de chaque page et facilitent la lecture, permettant par exemple aux écoliers de s'y retrouver pour préparer leurs exposés. Un glossaire et un index fournissent d'autres possibilités d'accès aux informations. C'est par ce côté fonctionnel que l'ouvrage se révèle un véritable atlas.

*Romain Pigeaud, Département de Préhistoire
du Muséum national d'histoire naturelle, Paris*



Archéologie

L'obsidienne

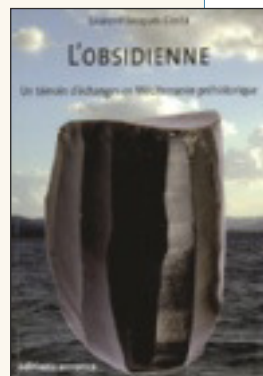
Un témoin d'échanges en Méditerranée préhistorique

Laurent-Jacques Costa
(112 pages, 16 euros), Errance, 2007.

L''obsidienne est un verre volcanique qui se trouve sous forme de coulées sur les pentes de certains volcans. Utilisée pour la première fois comme matière première il y a un million d'années en Éthiopie, elle fut surtout exploitée à partir du Néolithique, dans la vallée de l'Euphrate comme dans le bassin méditerranéen. Matériau très coupant, elle était recherchée, pour façonner les lames de haches, par exemple. D'une grande rareté, aux gisements néanmoins facilement identifiables, l'obsi-

dienne est d'un grand secours aux archéologues : ils peuvent, grâce à elle, identifier des réseaux d'échanges sur une grande échelle. Laurent-Jacques Costa dresse ici un vaste tableau des connaissances actuelles après plus de 40 ans de recherches sur le sujet. Puis il prend date pour de nouvelles études, fondées sur des hypothèses de travail qui replacent l'homme au centre de nos préoccupations. Car, selon L.-J. Costa, « il semble que, dans bien des cas, la description des vestiges se substitue au but même de leur étude, c'est-à-dire que l'on oublie qu'en sciences humaines, c'est l'homme qu'on est censé étudier ». Il fallait que cela soit dit.

R. Pigeaud



Darwin hérétique

Thomas Lepeltier (250 pages, 19 euros), Seuil, 2007.

L'Angleterre victorienne a accordé des funérailles solennelles à Darwin. C'est que, à l'époque, on ne voyait dans la théorie de l'évolution aucune attaque contre les fondements de la religion. Cette interprétation n'est venue qu'après. Le rapport entre la science et la religion évolue plus vite qu'on ne le croit, et le front de la guerre entre évolutionnistes et créationnistes ne cesse de se déplacer. Ce livre est remarquable par son aptitude à exposer les arguments et contre-arguments des uns et des autres, et à expliquer les raisons des divers protagonistes sans caricaturer leurs positions.



Cœur d'Afrique

Sous la dir. de Jean-Marie Hombert et Louis Perrois (225 pages, 35 euros), CNRS Éditions, 2007.



Dans cet ouvrage, une équipe de chercheurs français et gabonais entreprend de rendre justice à Paul Du Chaillu (1831-1903). À 17 ans, cet explorateur débarque au Gabon, où il fera trois séjours successifs. Son œuvre de naturaliste et ses observations ethnographiques sont méconnues, mais aussi importantes que celles de Livingstone ou de Brazza. En 16 passionnants chapitres, qui sont autant d'articles d'auteurs érudits, ce livre

fort bien illustré ressuscite le Gabon précolonial avec ses ethnies, ses nombreux animaux disparus et ses mystérieuses forêts denses.

Une histoire du temps et des horloges

Marie-Christine de la Souchère (170 pages, 16 euros), Ellipses, 2007.

Neuf chapitres augmentés d'appendices racontent de façon simple et claire l'histoire de la mesure du temps. L'évolution va des premiers étalons de temps, des premiers calendriers, des premières horloges pour en arriver à l'invention, puis au perfectionnement technique de la notion de temps universel. Efficace et idéal pour quiconque n'a pas le temps et veut se faire, en deux temps et trois mouvements, une culture sur le... temps !



La physique par les objets quotidiens

Cédric Ray et Jean-Claude Poizat (160 pages, 22,50 euros), Belin, 2007.



«Ceux qui utilisent négligemment les miracles de la science et de la technologie en ne les comprenant pas plus qu'une vache ne comprend la botanique des plantes qu'elle broute avec plaisir devraient tous avoir honte », a dit Einstein. C'est pourquoi de tels livres sont toujours bienvenus. De quoi est composé un écran LCD ? Quel est le secret de la géolocalisation GPS ? Sur quel effet repose la lecture des données enregistrées sur disque dur ? Et les fours à micro-ondes, les ampoules lumineuses,

les détecteurs de fumée ? Les objets de la vie quotidienne sont souvent de profonds mystères, même pour les physiciens et les ingénieurs. Cet ouvrage en éclaire 16 en autant de chapitres structurés par questions successives suivies de réponses illustrées d'images et de schémas.

Sciences de la Terre

Volcans

Olivier Grunewald et Jacques-Marie Bardintzeff (216 pages, 39,90 euros), Éditions du Chêne, 2007.

De grande qualité, les photographies qui constituent l'intérêt premier de ce livre sont pour la plupart récentes, certaines datant même de 2005 ou de 2006. Point fort : les points de vue adoptés par le photographe sont non seulement originaux, mais semblent étranges tant l'artiste a su jouer avec la lumière. Le choix de présenter les photographies par séries est particulièrement intéressant, car il donne une grande dynamique à l'ouvrage. L'univers des volcans y semble irréel, ce qui nous fait voyager dans notre imaginaire.

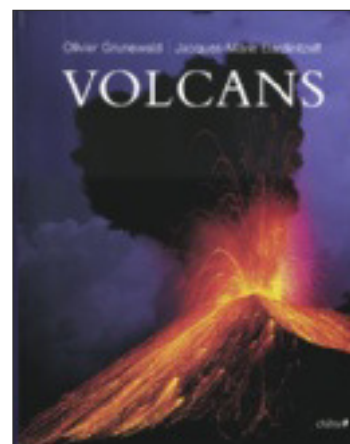
Dans le même temps, notre imaginaire est aussi invité à devenir parfois scientifique. L'introduction résume en quelques notions de base l'histoire de la Terre et comment les volcans se situent dans cette histoire, leur situation géographique et leur influence. La première partie donne les principales caractéristiques des éruptions, qui servent à distinguer trois thèmes photographiques : les « panaches explosifs », les « coulées brûlantes » et les « marmites du diable ». Cette organisation diffère de la présentation classique des types de volcans. Le texte court n'est pas nécessairement exhaustif et fait référence à quelques éruptions marquantes du siècle dernier, telles celles de la montagne Pelée en 1902 ou du Nevado del Ruiz en 1985, mais ce sont surtout les images qui prévalent ici et nous font voyager. Les séries de photographies font défiler des panaches, puis des coulées de lave et des nuées ardentes et enfin des lacs de lave, solidifiés ou non. Je mentionnerai ici deux clichés qui m'ont particulièrement marquée, l'un du sommet du Santiaguito, l'autre de la caldeira de l'Ert'a Ale avec la Lune en second plan.

La deuxième partie montre les paysages singuliers que seuls les volcans peuvent façonner. Les premières séries d'images concernent la morphologie plus ou moins conique des édifices, les deuxièmes les différents types de cratères, et les dernières, la matière volcanique elle-même avec des photographies de laves cordées et notamment un cliché de l'impressionnante chute d'Aldeyjarfoss, en Islande.

La troisième partie est consacrée aux sources et phénomènes géothermiques, et présente de superbes photographies des geysers du parc de Yellowstone, aux États-Unis.

Anne Le Friant,

Institut de physique du globe de Paris-CNRS



Espace

À la conquête de l'espace

De Spoutnik à l'homme sur Mars

Jacques Villain (315 pages, 25 euros), Vuibert, 2007.

Le 4 octobre 1957, la fusée soviétique *R7 Semiorokha* devenait le premier lanceur spatial en mettant *Spoutnik-1* sur orbite autour de la Terre. Vite célèbre, son bip-bip pouvait être capté par les radioamateurs de la planète entière. Quelques années plus tard, le 12 avril 1961, Youri Gagarine, à bord de *Vostok-1*, faisait un tour autour de la Terre en une heure et 48 minutes. En pleine guerre froide, la conquête de l'espace avait commencé. Ce livre nous la raconte, facette par facette.

Comme nous l'explique l'auteur, après *Spoutnik-1*, les satellites montreront très vite tout ce qu'ils peuvent apporter aux hommes. Leurs applications sont innombrables dans des domaines variés : la téléphonie, Internet, la télévision, la radiodiffusion numérique, la météorologie, la télé-médecine dans les zones déshéritées, etc. Le satellite s'est révélé être un facteur de progrès et de paix.

La conquête de l'espace, c'est aussi l'exploration de l'Univers, à l'aide de sondes automatiques qui vont visiter des lieux particulièrement hostiles, mais aussi à l'aide de télescopes tels que *Hubble*, qui a produit une moisson impressionnante de résultats. Même si les avancées acquises grâce aux satellites et aux sondes sont énormes, elles ont une image très technique si on les compare aux voyages effectués par l'homme dans l'espace. Pourtant, la présence d'hommes dans l'espace a peu d'intérêt, sauf par exemple quand il a fallu réparer le télescope *Hubble*. Du reste, la part du rêve s'est énormément réduite depuis l'époque des missions *Apollo*, épopée résumée par le célèbre « un petit pas pour un homme, un bond de géant pour l'humanité ». Alors, faudra-t-il que l'homme aille sur Mars pour que le rêve renaisse ?

Les événements qui, au cours des 50 dernières années, ont fait la conquête de l'espace, nous sont racontés par l'auteur à la faveur d'une présentation thématique. Il nous entraîne aussi sur quelques-unes des pistes dont le début est déjà connu, afin de tenter avec nous d'imaginer l'avenir. D'une lecture facile, son livre est bien documenté, très précis, facile à lire et rempli de détails et d'anecdotes. Il constitue une excellente lecture pour

se rendre compte à quel point la conquête spatiale est avant tout une aventure humaine. Et c'est justement ce côté si humain qui fait que les réalisations pour la rendre possible défient l'imaginaire.

Jean Cousteix, CERT-ONERA, Toulouse



Atmosphère, océan et climat

Robert Delmas, Serge Clauzy, Jean-Marc Verstraete et Hélène Ferré (288 pages, 29 euros), Belin, 2007.

La machine climatique se dérègle, nous assure-t-on presque quotidiennement. Mais en quoi consiste donc cette machine ? Ce livre répond en passant en revue l'ensemble de ses « organes », qu'il s'agisse de l'atmosphère, de ses nuages, des océans et de ses mouvements. L'information est apportée sous forme de textes et de nombreux

schémas en couleur. Précis et simple.

Le Léman et sa vie microscopique

J.-C. Druart et G. Balvay (180 pages, 45 euros), Quae, 2007.

Destiné aux hydrobiologistes, cet ouvrage intéressera également ceux qui se passionnent pour l'histoire et la géologie du Léman, ainsi que l'évolution de la qualité de ses eaux. Après une présentation du bassin-versant du Léman, les auteurs exposent et analysent son réseau trophique (la chaîne alimentaire) et toutes les sortes de plancton vivant dans le grand lac alpin. De belles photographies identifient les nombreux organismes mentionnés. Courbes et illustrations soutiennent le texte. Un outil de spécialiste, mais à la portée de tous.



Atlas des pôles

Éric Canobbio et Aurélie Boissière (80 pages, 15 euros), Autrement, 2007.

Sanctuarisé depuis 1959, l'Antarctique est un continent vide, tandis que l'Arctique est avant tout un océan entouré de terres. Quels sont les enjeux concernant les deux extrémités de notre planète ? Le mieux est de le dire en cartes : dans cet ouvrage, une multitude de cartes illustrent de brefs et efficaces articles traitant du relief, du climat, de la biodiversité, de la pêche, de la guerre froide, des populations autochtones, du développement durable des déserts glacés, du statut juridique des régions polaires, etc.



Ventoux, géant de nature

Guide des richesses biologiques du mont Ventoux

Frédéric Melki et Maxime Briola (250 pages, 39 euros), Biotopie, 2007.

Œuvre d'un groupe de naturalistes passionnés, cet ouvrage présente les paysages, la géologie, les fossiles et la flore de cette montagne du Vaucluse. Leur passion s'exprime notamment par les centaines de magnifiques photographies qu'ils ont accumulées au cours de leurs randonnées. Les textes regorgent de termes spécifiques à la botanique, à la géologie, à la région... tout en restant bien compréhensibles. Leur remarquable clarté est renforcée par un intelligent système de paragraphes et de mises en exergue des termes évoqués.



Logique

Les démons de Gödel

Pierre Cassou-Noguès (280 pages, 21 euros),
Seuil, 2007.

Un petit miracle, ce livre ! L'auteur, agrégé de mathématiques, docteur en philosophie, écrit une langue claire, sans jargon. Certes, le livre exige de l'attention – il s'agit de la pensée de Gödel, tout de même ! – mais son auteur a fait un remarquable effort de lisibilité. Parfois, même, il se met en scène dans de brèves fictions, où il ne cache pas les quasi-hallucinations que finit par susciter en lui la lecture des brouillons laissés par Gödel.

L'auteur a eu accès aux milliers de pages que la veuve de Gödel a léguées à l'*Institute for Advanced Study* de Princeton. Leur dépouillement, inachevé, n'est pas facile. Gödel écrivait en anglais, en allemand... et en Gabelsberger – une forme de sténographie qu'apprenaient les élèves germanophones du début du xx^e siècle. Aujourd'hui, seuls de rares spécialistes peuvent la déchiffrer. Gödel a peu publié, après son théorème d'incomplétude. Il cachait même ses notes, de peur qu'on le prenne pour fou.

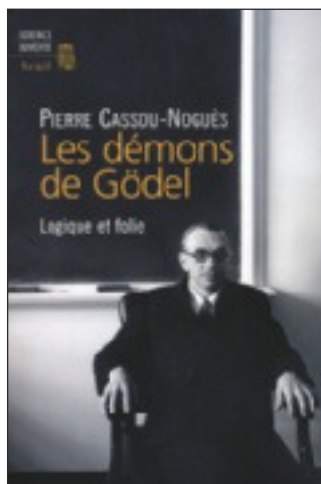
Justement, il ne fait guère de doute aux yeux de P. Cassou-Noguès que Gödel était ce que le sens commun appelle

« fou ». Pour Gödel, tout a un sens dans l'Univers, il n'est pas de coïncidence qui ne doive être expliquée. Le monde mathématique est peuplé d'anges, êtres qui accompagnent les objets mathématiques et sont parmi les idées comme les hommes sont dans la matière. Une question inquiétante se pose alors : y a-t-il des liens entre la « folie » de Gödel et son œuvre logique ? La réponse de l'auteur, inquiétante aussi, est oui. Non pas que le théorème d'incomplétude soit « fou » – évidemment pas ! – mais il s'insère bien dans le système de pensée de Gödel. Pour Gödel, du moment que la raison peut se poser quelque problème, elle peut aussi résoudre celui-ci. Il faut donc que les hommes soient capables de résoudre tous les problèmes posés en arithmétique. Or le cerveau, ne pouvant connaître qu'un nombre fini d'états physiques différents, est une machine de Turing. Mais, d'après le théorème d'incomplétude, une machine de Turing ne peut pas résoudre tous les problèmes diophantiens. Donc l'esprit humain ne se réduit pas à son cerveau !

Autre aspect du thème « logique et folie », la raison de l'existence du temps. Selon Gödel, le monde est contradictoire et la seule façon pour Dieu de permettre aux hommes de penser ces contradictions a été de les développer

dans le temps. Conception étrange ? On en trouvera bien d'autres exposées dans ce livre.

Didier Nordon, Université de Bordeaux



POUR LA SCIENCE

Directrice de la rédaction - Rédactrice en chef : Françoise Pétry

Pour la Science :

Rédacteurs en chef adjoints : Maurice Mashaal, Loïc Mangin

Rédacteurs : François Savatier, Philippe Ribeau-Gésippe,

Bénédicte Salthun-Lassalle

Dossiers Pour la Science :

Rédactrice en chef adjointe : Bénédicte Leclercq

Génies de la Science :

Rédactrices : Marie-Neige Cordonnier et Bénédicte Leclercq

Cerveau & Psycho :

Rédacteurs : Sébastien Bohler et Bénédicte Salthun-Lassalle

Directrice artistique : Céline Lapert

Secrétariat de rédaction/Maquette : Annie Tacquenot,

Sylvie Sobelman, Pauline Bilbault, Raphaël Queruel, Ingrid Leroy

Site Internet : David Martin

Marketing et Publicité : Philippe Rolland, assisté de Marina Ballini

Direction financière : Anne Gusdorf

Direction du personnel : Marc Laumet

Fabrication : Jérôme Jalabert, assisté de Marianne Sigogne

Presse et communication : Susan Mackie

Directeur de la publication et Gérant : Marie-Claude Brossollet

Conseillers scientifiques : Philippe Boulanger et Hervé This

Ont également participé à ce numéro : Isabelle Allemand,

Robert Barbault, Bettina Debü, Raphaël Ilan, Christophe Pichon,

Geneviève Rougon, Daniel Tacquenot

SERVICE ABONNEMENTS

Ginette Grémillon : Tél. : 01 55 42 84 04

Abonnements pour la Belgique : EDIGROUP Belgique Sprl - 5 place du Champs

de Mars - 20th floor - 1050 Bruxelles - tel 070/233 304 - abobelgique@edigroup.org

Abonnements pour la Suisse : EDIGROUP SA - 39 rue Peillonex CH 1225

Chene Bourg - Tel 022/860 84 01 - abonne@edigroup.ch

8 rue Férou, 75278 PARIS CEDEX 06 • Tél : 01-55-42-84-00 • www.pourlascience.com
Commande de dossiers ou de magazines : 08-92-68-11-40 (de l'étranger, (+33) 2 37 82 06 62)

PUBLICITÉ France

Directeur de la Publicité : Jean-François Guillotin (jf.guillotin@pourlascience.fr),

assisté de Nada Mellouk

Tél. : 01 55 42 84 28 ou 01 55 42 84 97 • Télécopieur : 01 43 25 18 29

DIFFUSION DE POUR LA SCIENCE

Canada : Edipresse : 945, avenue Beaumont, Montréal,

Québec, H3N 1W3 Canada.

Suisse : Servidis : Chemin des châlets, 1979 Chavannes - 2 - Bogis

Belgique : La Caravelle : 303, rue du Pré-aux-oies - 1130 Bruxelles.

Autres pays : Éditions Belin : 8, rue Férou - 75278 Paris Cedex 06.

SCIENTIFIC AMERICAN Editor in chief : John Rennie. Board of editors : Mariette Di Christina, Ricky Rusting, Philip Yam, Gary Stix, Michelle Press, Wayt Gibbs, Mark Alpert, Steven Ashley, Graham Collins, Mark Fischetti, Steve Mirsky, George Musser, Christine Soares. Chairman : Brian Napack. Vice president and managing director international : Dean Sanderson. Vice president : Frances Newburg. Chairman Emeritus : John Hanley.

Toutes demandes d'autorisation de reproduire, pour le public français ou francophone, les textes, les photos, les dessins ou les documents contenus dans la revue « Pour la Science », dans la revue « Scientific American », dans les livres édités par « Pour la Science » doivent être adressées par écrit à « Pour la Science S.A.R.L. », 8, rue Férou, 75278 Paris Cedex 06.

© Pour la Science S.A.R.L.

Tous droits de reproduction, de traduction, d'adaptation et de représentation réservés pour tous les pays. La marque et le nom commercial « Scientific American » sont la propriété de Scientific American, Inc. Licence accordée à « Pour la Science S.A.R.L. ».

En application de la loi du 11 mars 1957, il est interdit de reproduire intégralement ou partiellement la présente revue sans autorisation de l'éditeur ou du Centre français de l'exploitation du droit de copie (20, rue des Grands-Augustins - 75006 Paris).